



Universiteit
Leiden
The Netherlands

A random matrix perspective of cultural structure: groups or redundancies?

Babeanu, A.I.

Citation

Babeanu, A. I. (2021). A random matrix perspective of cultural structure: groups or redundancies? *Journal Of Physics: Complexity*, 2(2). doi:10.1088/2632-072X/abc859

Version: Publisher's Version

License: [Creative Commons CC BY 4.0 license](https://creativecommons.org/licenses/by/4.0/)

Downloaded from: <https://hdl.handle.net/1887/3278064>

Note: To cite this publication please use the final published version (if applicable).

PAPER • OPEN ACCESS

A random matrix perspective of cultural structure: groups or redundancies?

To cite this article: Alexandru-Ionu Bbeanu 2021 *J. Phys. Complex.* **2** 025008

View the [article online](#) for updates and enhancements.

You may also like

- [Effects of Non-Uniform Temperature on in-Situ Current Distribution and Non-Uniform State of Charge Measurements for \$\text{LiFePO}_4\$ and \$\text{LiNiMnCoO}_2\$ Cells](#)
Matthew Paul Klein and Jäe Wan Park
- [Thomson scientists avoid redundancies but lose work](#)
Bradford Smith
- [Lorentz- and permutation-invariants of particles](#)
Ben Gripaios, Ward Haddadin and Christopher G Lester

OPEN ACCESS



CrossMark

RECEIVED

28 July 2020

REVISED

2 November 2020

ACCEPTED FOR PUBLICATION

6 November 2020

PUBLISHED

29 January 2021

Original content from this work may be used under the terms of the [Creative Commons Attribution 4.0 licence](#).

Any further distribution of this work must maintain attribution to the author(s) and the title of the work, journal citation and DOI.



PAPER

A random matrix perspective of cultural structure: groups or redundancies?

Alexandru-Ionuț Băbeanu

Lorentz Institute for Theoretical Physics, Leiden University, The Netherlands

Rotterdam School of Management, Erasmus University, The Netherlands

* Author to whom any correspondence should be addressed.

E-mail: a.i.babeanu@gmail.com

Keywords: cultural structure, random matrix theory, similarity eigenmode analysis, symbolic sequences, internal group uniformity, spin systems, community detection

Abstract

Recent studies have highlighted interesting properties of empirical cultural states—collections of cultural trait sequences of real individuals. Matrices of similarity between individuals may be constructed from these states, allowing for more insights to be gained using random matrix techniques, approach first exploited in this study. We propose a null model that enforces, on average, the empirical occurrence frequency of each possible trait. With respect to this null model, the empirical matrices show deviating eigenvalues, which may be signatures of subtle cultural groups. However, they can conceivably also be artifacts of arbitrary redundancies between cultural variables. We study this possibility in a highly simplified setting, allowing for a side-by-side mathematical comparison of the two scenarios (groups and redundancies). The scenarios are shown to be completely indistinguishable in terms of deviating eigenvalues, confirming that the latter can in general be signatures of either redundancies or groups. The scenarios can be distinguished after evaluating the eigenvector uniformities and the associated deviations from null model expectations. This provides a uniformity-based validation criterion, which is reliable when searching for groups that are internally uniform, but fails when these exhibit significant internal non-uniformity. For empirical data, all the relevant eigenvector uniformities are compatible with the null model, indicating the absence of any internally uniform groups. Although there are various indications that some of the deviating eigenvalues could correspond to internally non-uniform groups, a generic procedure for distinguishing such groups from redundancy artifacts requires further research.

1. Introduction

Understanding the complex behavior of social systems greatly benefits from constructively combining the increasing amount of empirical data with a variety of quantitative, theoretical approaches, often originating in the natural sciences [1, 2]. Although much of this interdisciplinary research focuses on the network and connectivity aspects of social systems [3], efforts are also being made for understanding a complementary aspect: the formation and dynamics of opinions, preferences, attitudes and beliefs, more generically referred to as ‘cultural traits’ [4]. In particular, recent studies have placed a stronger emphasis on using empirical data about the cultural traits of real individuals [5–9]. Such data is typically recorded within a short period of time from a random sample of people in a population, via a social survey with a large number of questions/items, so that a vector (or sequence) of cultural traits can be constructed for every individual, where each trait is an answer to one of the questions. The collection of all cultural vectors constructed from one empirical source is called an empirical ‘cultural state’, or an empirical ‘set of cultural vectors’, since it can be used to empirically specify the initial conditions of an Axelrod-type model of cultural dynamics [10]. Using previously developed tools [5, 6] that relied on models of cultural and opinion dynamics, reference [7] showed that empirical cultural states are characterized by properties that are highly robust across different datasets. These properties have

been further explored [8, 9] but not entirely understood. The generic empirical structure appears to be largely captured by the matrix of cultural similarities between individuals, which are computed for all pairs of cultural vectors. This opens the possibility to further investigate the empirical structure by means of a random matrix approach.

Random matrix theory [11, 12] has been successfully used for a variety of applications, among which the analysis of financial systems [13] is an important highlight, with a traditional focus on stock markets [14] and recent extensions to credit default swap markets [15] and the global banking network [16]. It also remains an active area of theoretical research [17, 18]. The framework deals with various properties of random matrices, under certain distributional assumptions. The associated statistical ensemble of matrices is used to compute the expected values (or even the probability distributions) of interesting, matrix dependent quantities. These theoretical expectations can be compared to empirical counterparts evaluated on matrices that encode information about the real world systems that are being studied. Statistically significant deviations of the empirical quantities are then interpreted as interesting, non-trivial structural properties of the respective systems. The statistical ensemble thus acts as a null model or benchmark, with respect to which interesting empirical structure is to be identified. The focus is on the eigenvalue spectrum of the empirical matrix, which very often consists of correlations between time series associated, for instance, to stock dynamics [14] or neural activity, at the level of single cells [19] or brain regions [20]. In such cases, the appropriate assumptions of randomness are largely captured by the Marchenko–Pastur [21] law, which gives a limiting distribution for the spectrum. The empirical eigenmodes whose eigenvalues are larger than what is expected based on the Marchenko–Pastur law are interpreted as joint dynamical patterns in terms of which the non-trivial behavior of the system can be understood, while the other are interpreted as noise components. Recently, reference [22] extended this approach to similarity matrices constructed from categorical data, where each entry of the matrix is a similarity between two time series of discrete symbols. For instance, for one of the datasets in reference [22], each sequence of symbols corresponds to an electoral constituency of India, with different symbols associated to different winning parties and successive time steps associated to successive elections.

This study extends the approach to spectra of empirical matrices of cultural similarities, constructed from data previously used in references [5–9]. Instead of relying on analytic procedures for estimating and filtering the noise, numerical methods are extensively used, largely based on a null model that is first introduced here, namely ‘restricted randomness’ (or *r*-random), which is shown to be very appropriate in this context. This allows us to evaluate the distribution of the upper boundary of the noise region (‘the bulk’). We show in section 2 that there are several empirical eigenvalues significantly above this boundary, for each empirical matrix. These ‘deviating eigenmodes’ capture the structure of empirical data, since they are incompatible with the null hypothesis behind restricted randomness. Hence, this manuscript will often refer to them as ‘structural modes’.

It is tempting to interpret the structural modes as manifestations of cultural groups, in a manner similar to time series analysis [14], suggesting that individuals fall under several classes, categories or clusters that are inherent to the system. This is particularly intriguing, given that reference [8] provides indirect evidence for cultural structure being governed by a small number of cultural prototypes potentially induced by universal ‘rationalities’. However, it is important to keep in mind that the empirical data also shows pairwise correlations between cultural variables (or ‘features’), that are at least partly due to arbitrary, dataset-dependent redundancies (semantic overlaps) between the corresponding survey items, as previously pointed out [5–7]. Since these correlations are not retained by restricted randomness, it is possible that deviating eigenmodes are a direct consequence of arbitrary redundancies.

The question of whether deviating eigenmodes are signatures of authentic groups of individuals or just artifacts of arbitrary redundancies between variables motivates the rest of the study. First, section 3 shows that it is mathematically meaningful to differentiate, at least in principle, between a ‘redundancies scenario’ and a ‘groups scenarios’. This is done by constructing for each scenario a probabilistic (cultural state generating) toy model that captures its essence, while operating in a common, highly simplified setting, in isolation from empirical data. While remaining in this setting, section 4 shows that in order to distinguish between the two structural scenarios, deviating eigenvalues are not enough: a measure of eigenvector uniformity needs to be used in a complementary manner; structural modes corresponding to authentic groups then exhibit significantly higher uniformity than predicted by the null model. Section 5 makes use of these insights in the empirical setting, while presenting an enhanced analysis of the cultural states studied in section 2: the evaluation of empirical eigenvector uniformities suggests that all structural modes are redundancy artifacts. However, section 6 provides indications that at least some of the structural modes are actually associated to authentic groups, but these groups are in a certain sense too subtle to be recognized by the uniformity-based approach. Specifically, the eigenvector uniformity criterion works for underlying groups that are internally uniform and have a hard boundary, but not for groups that are internally non-uniform and have a soft boundary; real-world

cultural groups are expected to be of the latter type. Section 7 discusses the implications of this study and future research prospects, while section 8 concludes the manuscript.

2. Eigenvalue spectra for empirical data and null model

In this section, the eigenvalue spectra of empirical matrices of cultural similarities are evaluated. These are compared to eigenvalue spectra and distributions obtained from the ‘restricted random’ null model, which is first introduced here, and chosen as a good benchmark with respect to which interesting structure is to be measured. Similarity matrices are numerically generated by randomly sampling from the statistical ensemble of restricted randomness, which enforces empirical information that is not of interest—information that, on *a priori* grounds, clearly has more to do with arbitrary survey design choices than with any authentic cultural structure. The details behind the choice of the null model are elaborated in appendix A. Before presenting the results, some mathematical clarifications are given with respect to the computation of similarity matrices and the spectral decomposition procedure.

A cultural similarity matrix is a square, $N \times N$ matrix obtained from N cultural vectors, which are all defined with respect the same set of F cultural features (variables or dimensions). Each feature can take one of q^k possible, discrete values, called ‘cultural traits’, where k labels the features, according to some order that is arbitrary, but consistent across all vectors. Moreover, each feature can be either nominal, marked as $f_{\text{nom}}^k = 1$, or ordinal, marked as $f_{\text{nom}}^k = 0$, which affects how its similarity contribution is defined. Each entry s_{ij} of the similarity matrix is then computed according to:

$$s_{ij} = \frac{1}{F} \sum_{k=1}^F \left[f_{\text{nom}}^k \delta(x_i^k, x_j^k) + (1 - f_{\text{nom}}^k) \left(1 - \frac{|x_i^k - x_j^k|}{q^k - 1} \right) \right], \quad (1)$$

encoding the similarity between vectors i and j , where δ stands for the Kronecker delta function and x_i^k and x_j^k denote the traits recorded with respect to feature k in vectors i and j respectively—for the ordinal case, it is important that x_i^k and x_j^k take discrete, rational values between 1 and q^k , while for the nominal case they only need to take symbolic values from any (feature-specific) alphabet. Note that the similarity measure in equation (1) is an arithmetic average of the similarity contributions of the F cultural features, in agreement with references [5–9]—although in these studies most concepts are presented in terms of cultural distances d_{ij} , these have a trivial relationship to cultural similarities: $d_{ij} = 1 - s_{ij}$. For an empirical matrix, each vector i corresponds to one individual in the real world, each feature k to one question or item in the questionnaire used to collect the data, so that the realized trait x_i^k , which lies at the intersection between vector i and feature k , corresponds to the answer/rating given by individual i to question/item k . For a matrix generated based on a null model, the N vectors are generated according to the specified random procedure, while retaining (at least) the empirical data format, namely the type f_{nom}^k and range q^k of each feature k . Note that, in contrast to the empirical symbolic sequences used in reference [22], cultural vectors have no axis of time, so everything is equivalent up to a reordering of the cultural features, as long as this is done consistently for all cultural vectors. This is irrelevant for any of the mathematical operations involved in the analysis here, but is relevant for the interpretation: cultural vectors capture no time-evolution, and should be interpreted as instantaneous, multidimensional opinion profiles, rather than as dynamical, one-dimensional profiles.

Equation (1) implies that such a similarity matrix is real and symmetric. Thus, according to the spectral theorem, it has N real eigenvalues with N associated orthonormal eigenvectors with real entries, so the matrix can be decomposed in the following way:

$$s_{ij} = \sum_{l=1}^N \lambda_l v_l^i v_l^j, \quad (2)$$

where ‘ λ_l ’ and ‘ v_l ’ are used to denote the l th highest eigenvalue and, respectively, the eigenvector associated to it, while v_l^i is the i th entry of eigenvector v_l . Throughout this study, special attention is paid to λ_1 and λ_2 , the highest and second highest eigenvalues of various similarity matrices, also denoted as the ‘leading’ and ‘subleading’ eigenvalues respectively. In parallel, ‘ λ ’ is used to denote any generic eigenvalue. More notation will be introduced below, as needed.

It is remarkable that all eigenvalues of any similarity matrix computed according to equation (1) are bound to the positive real axis: $\lambda_l \geq 0, \forall l$. Thus, in this study, no eigenvalue-related figure axis needs to be concerned with values smaller than 0.0. This positive semidefiniteness property is rigorously shown to hold in appendix B. The property may be relevant beyond the analysis of cultural states, since it holds for any similarity matrix belonging to what one may call the ‘Hamming–Manhattan’ class: a matrix whose elements s_{ij} satisfy $s_{ij} = 1 - d_{ij}$, where every d_{ij} is a combined, Hamming–Manhattan distance, with nominal and ordinal features corresponding to the Hamming and Manhattan contributions respectively.

Table 1. Empirical data format summary. The table shows the total number of features F , as well as how many of them (m) have a certain combination of type (f_{nom}) and range (q), where f_{nom} indicates whether a feature is nominal (1) or ordinal (0), while q indicates the number of associated traits. This is shown for each of the three datasets used in this study: EBM, GSS and JS. Note that the sum of all m values is equal to F for each dataset.

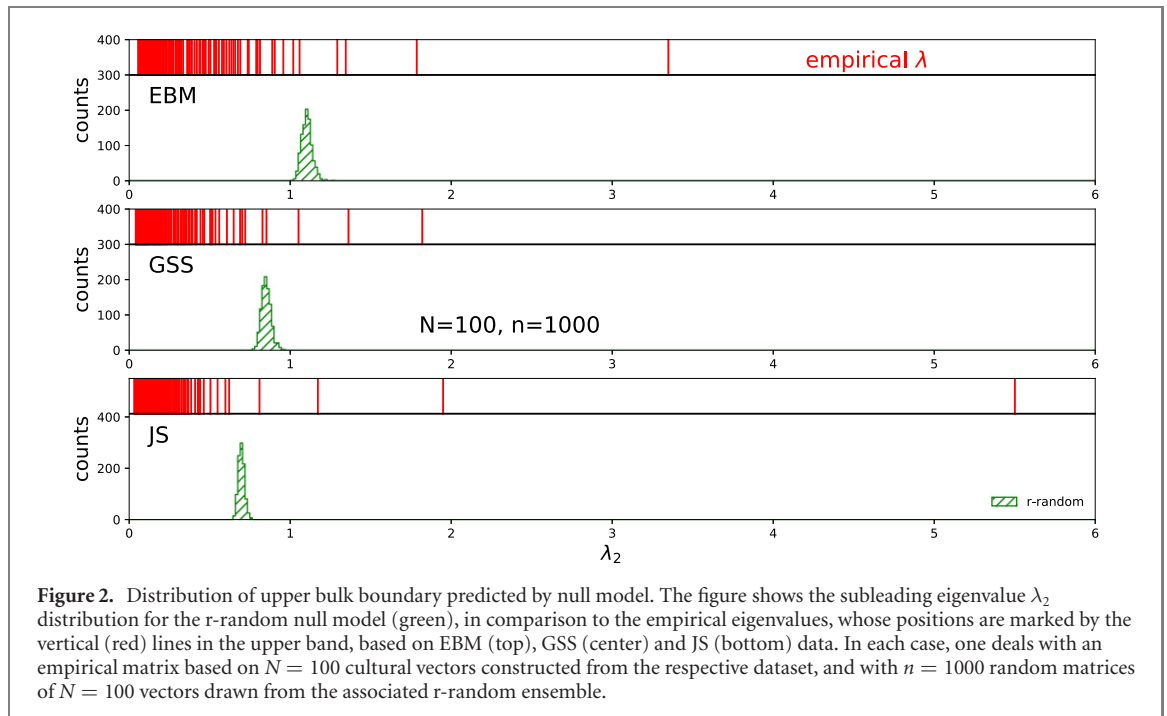
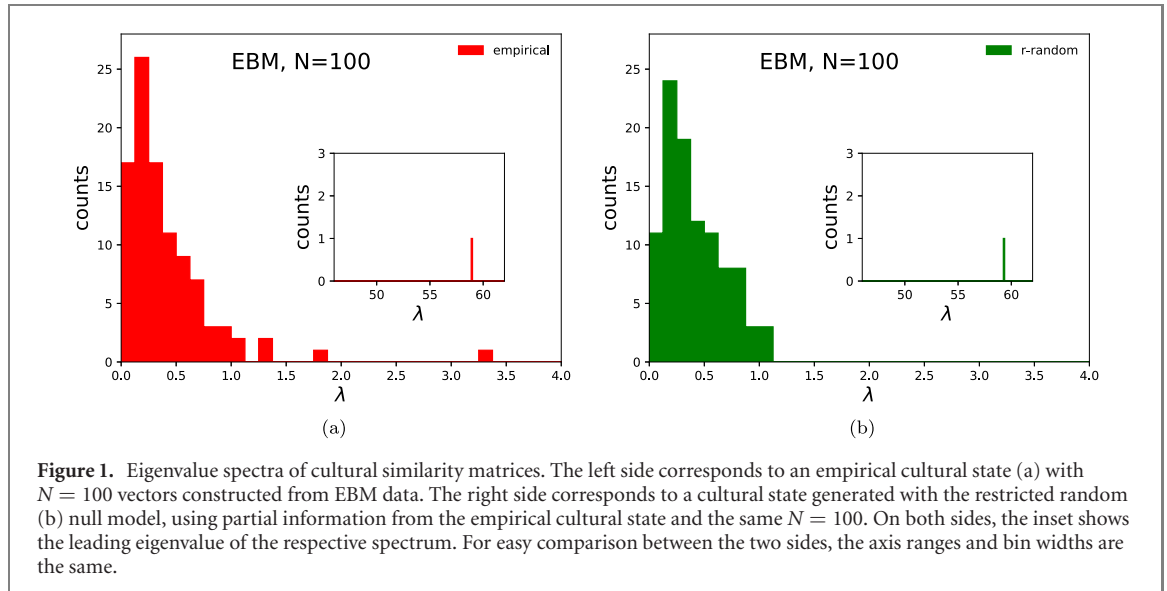
Dataset	EBM	GSS	JS
F	144	122	128
(m, f_{nom}, q)	(40, 1, 3) (38, 0, 3) (34, 0, 5) (9, 0, 7) (8, 0, 4) (7, 1, 4) (5, 1, 5) (1, 1, 16) (1, 1, 11) (1, 0, 6)	(67, 0, 5) (26, 0, 4) (18, 0, 3) (8, 1, 3) (1, 0, 9) (1, 0, 7) (1, 0, 6)	(128, 0, 7)

All similarity matrices used in this study are based on sets of $N = 100$ cultural vectors, regardless of whether they are empirical, generated with a null model, or generated with one of the toy models introduced in section 3 and section 4. Moreover, all these matrices satisfy $s_{ii} = 1.0, \forall i$, as a consequence of equation (1), meaning that the trace is always: $\sum_i s_{ii} = N = 100$. Since diagonalization preserves the trace, the eigenvalues are also bound to add up to $\sum_l \lambda_l = N = 100$ for every matrix.

Three datasets are used in this study (also used in references [7–9]) for constructing empirical cultural states: first, an edition of the Eurobarometer [23] (EBM), which records attitudes and opinions of European Union citizens on various topics concerning technology, the environment and other policy issues (mostly pertaining to European integration); second, an edition of the General Social Surveys [24] (GSS), which records attitudes and opinions of United States citizens on a variety of topics; third, the Jester [25] (JS) 2 data, recording ratings of jokes. The data collection is based on face-to-face interviews for EBM and GSS, and on an interactive web application for JS; in each case, the raw data is available online—provided for free by its owner. For the purpose of this study, the data is formatted according to the procedure described in reference [7], which leads to a certain combination of cultural features for each dataset, as summarized by table 1. For EBM and GSS data, the formatting involves a multitude of operations, many of which require careful human inspection. For instance, one needs to eliminate survey items that are either not subjective enough or otherwise have other problems that make them inappropriate as cultural features—such as not being presented to all respondents or being defined in a manner that explicitly depends on the respondent’s identity. Also, one needs to consistently merge different ballots/forms—subsets of respondents for which some of the survey items are formulated in slightly different ways, while appearing as duplicates in the raw data—as well as gracefully deal with neutral/missing answers. For JS, the formatting is simpler, mostly involving the conversion of quasi-continuous ratings between -10 and 10 to a discrete, ordinal scale with $q = 7$ traits, while also eliminating people that have not rated all the jokes.

Figure 1(a) shows the eigenvalue spectrum of an empirical similarity matrix computed based on $N = 100$ cultural vectors extracted from EBM data. The vertical axis gives the number of eigenvalues occurring in each bin along the horizontal axis. The inset focuses on the higher λ region of the horizontal axis, where the leading eigenvalue λ_1 is located. The high value of λ_1 is expected based on purely mathematical grounds [22], due to the overall positivity of any such similarity matrix. In most cases, all entries of the eigenvector associated to λ_1 have the same sign and very similar absolute values, meaning that, according to equation (2), the $\lambda_1 v_1^i v_1^j$ captures a large, highly uniform, positive component of the matrix entries s_{ij} . The λ_1 eigenmode thus accounts for the overall tendency toward similarity of the entire system: its existence is a consequence of how similarity is defined, but its strength is controlled (see below) by the feature-level probability distributions (the trait frequencies). For this reason, the λ_1 mode will also be referred to as the ‘global mode’, like in reference [14]—in the context of correlation-based analysis of time-series, a global mode may or may not be present, depending on the nature and origin of the data.

Using the same format as figure 1(a), figure 1(b) shows the spectrum of a similarity matrix generated via ‘restricted randomness’ (abbreviated as ‘r-random’). Specifically, for every vector, each trait is chosen independently at random from the traits available at the level of the respective feature, with different probabilities attached to the possible traits, these probabilities being directly proportional to the empirical occurrence frequencies of the respective traits. Thus, in addition to retaining the empirical data format, restricted randomness



reproduces, on average, the empirical trait frequencies, whose values mostly reflect survey design choices. Conceptually, restricted randomness combines desirable properties of two other null models, based on uniform randomness and trait shuffling respectively, which have previously been used [5–9] for the analysis of cultural states, independently from eigendecomposition and random matrix notions—see appendix A.

The rough shape of the eigenvalue histogram is quite similar in the two panels of figure 1, meaning that empirical data contains a large amount of noise, which can be described well by restricted randomness. The match between the empirical and the r-random states is particularly remarkable at the level of the leading eigenvalue λ_1 , which, as elaborated in appendix A, constitutes an important validation of restricted randomness as a suitable null model.

Very important are the empirical outliers in figure 1(a), which encode empirical structure that is independent of feature-level non-uniformities, since they do not have correspondents in figure 1(b). In particular, these eigenvalues appear significantly higher than the random bulk region predicted by restricted randomness, with a lower boundary of $\lambda_- = 0.0$ and an upper boundary of $\lambda_+ \approx 1.1$. However, not much can be understood about the uncertainty around λ_+ based on figure 1 alone, since figure 1(b) uses only one matrix sampled from the r-random ensemble, which does not provide information about the flexibility of the r-random spectrum

under resampling. This is essential for deciding how many of the empirical outliers are to be regarded as structurally relevant rather than as noise components, a question which is particularly important for the lower λ outliers.

These limitations are addressed by the top part of figure 2, by showing the subleading eigenvalue distribution for restricted randomness. For comparison, the empirical eigenvalues are shown by the vertical (red) lines in the upper band. The λ_2 distribution is produced numerically by sampling $n = 1000$ sets of cultural vectors from the r-random statistical ensemble. With respect to this distribution, all four empirical outliers noted in figure 1(a) appear statistically significant, with a departure of at least two standard deviations from the mean. Figure 2 also replicates the analysis with empirical cultural states constructed from GSS and JS data. Both eigenvalue spectra show outliers that are significantly larger than what is expected based on the r-random null model: three such outliers are present for GSS and four for JS. The deviating eigenvalues are, on average, larger for JS than for EBM, and higher for EBM than for GSS.

Based on the results above, one can say that the empirical structure captured by matrices of cultural similarity is generally recognizable via eigenvalues that are significantly larger than what is expected based on a null hypothesis accounting for empirical trait frequencies: they are significantly higher than the subleading eigenvalue and much lower than the leading eigenvalue expected from this null hypothesis. For the rest of this study, the eigenpairs (eigenvector–value pairs) associated to these deviating eigenvalues will often be referred to as ‘structural modes’.

3. Two interpretations of structural modes

This section presents two possible interpretations for the structural modes of culture highlighted in section 2. On one hand, structural modes may be the effect of redundancies between cultural features, thus only retaining information about how the associated questions/items are constructed. On the other hand, structural modes may be an effect of genuine groups or grouping tendencies among the individuals, thus retaining information about the social system from which the data is extracted. We start by presenting the basic reasoning and intuition on which either interpretation is based. We then proceed to precise probabilistic formulations of the two scenarios, in a very simplistic setting: the redundancy scenario is realized as the ‘fully-connected Ising’ (FCI) model in section 3.1, while the groups scenario is realized as the ‘symmetric two-groups’ (S2G) model in section 3.2. Finally, in section 3.3, the mathematical properties of the two models are studied in order to check that they behave as expected and to better emphasize their differences.

The interpretation based on groups largely stems from basic linear algebra considerations, which are presented in appendix C. When combining these considerations with the findings in section 2, one concludes that structural modes are normalized linear combinations of the individuals, orthogonal to each other and to the global mode, with the highest possible self-similarities, of which the lowest is significantly higher than what is expected from restricted randomness. Each structural mode could thus indicate the presence of a group of highly similar individuals (in the context of time-series analysis, structural modes are often called ‘group modes’ [14]), or, perhaps more precisely, the presence of a tendency toward clustering around a certain locus in cultural space. Although it is not clear how a linear combination of individuals (or of cultural vectors) should be expressed in terms of cultural traits and features, this is not important for this study and does not affect the argument.

The interpretation based on redundancies comes from realizing that social surveys are imperfect, in the sense that one cannot guarantee the absence of semantic overlaps (redundancies or similarities) between the variables that are used. These lead to correlations between cultural features, which have been noticed in previous studies [5–7] and which are specific to the design of each dataset. It is conceivable that feature redundancies, if strong enough, could induce artifactual structural modes themselves. For example, if a large fraction of the associated items or questions are designed such that they exhibit strong semantic overlaps with each other, the similarity between individuals responding to any of these items in a certain way will be high, since these individuals will likely respond to all the other items in the same way. It appears likely that this situation would induce a structural mode. If this is the mechanism behind the structural modes shown in section 2, they do not provide information about the inherent organization of real-world culture, but just about the design of the ‘instrument’ used to ‘measure’ culture. Although redundancies between features would manifest as correlations between those features, authentic groups can also induce such correlations, so this aspect cannot be directly exploited for differentiating redundancies from groups.

There are thus fundamental reasons that make it very difficult to understand the extent to which structural modes of culture are due to details of the experimental setting and the extent to which they are due to authentic properties of the underlying system. Here, a first step in this direction is made, by formulating the two scenarios as mathematical, probabilistic models capable of generating (sets of) cultural vectors that are governed either by a pairwise coupling between cultural features (section 3.1) or by a grouping tendency (section 3.2). These

models are designed to work without any empirical input, in the same, simplest conceivable setting, consisting of F binary features—it does not matter whether these features are regarded as ordinal or nominal, since the two types of similarity contributions are equivalent if there are only $q = 2$ traits available, as can be seen from equation (1). For each feature, the two traits are marked as ‘−1’ and ‘+1’—although the former should be mapped to ‘0’ when computing similarities between vectors, if features are assumed ordinal. Each of the two models defines a statistical ensemble (and an associated cultural space distribution, in the language of references [7, 8]), according to which cultural vectors can be independently drawn, in a random, but non-uniform way. For both statistical ensembles, each feature-level probability distribution is uniform—the two traits have an equal probability of 0.5 attached. Note that, although both models are probabilistic in nature, neither of them is intended as a null model, since neither makes use of information from empirical data nor is it intended for direct, quantitative comparisons to empirical data, nor to be realistic to any extent. They are toy-models, used for illustrating conceptual differences and ambiguities between redundancies and groups in the context of cultural states. Nonetheless, they do provide an arena for studying and developing certain mathematical tools in a highly controlled setting, tools that can later be used for studying empirical data.

3.1. The first scenario: redundancies

This section explains the FCI model, in the context of generating (sets of) cultural vectors in a stochastic way. The purpose of this probabilistic model is to enforce a certain level of redundancy for all pairs of cultural features, controllable via one parameter, but as little as possible in addition. This can be done by properly choosing the probability distribution p taking as support the set of possible cultural vectors with F binary features, or, in other words, the set of possible spin configurations $\vec{S}$ with F lattice sites. Note that the support of this distribution has 2^F elements, which is the number of sites/points of the ‘cultural space’, according to the formalism in reference [7].

One needs to choose the maximally-random (thus minimally biased) probability distribution p that entails a certain level of feature–feature couplings. This is found by maximizing the Shannon entropy (equation (D1)) subject to two constraints: one enforcing the normalization of the probability distribution (equation (D2)), the other enforcing the overall level of pairwise coupling between cultural features (equation (D3)). This procedure is a realization of maximum-entropy inference introduced in reference [26], and is described in detail in appendix D. The resulting probability distribution can be expressed as:

$$p(\mu, F, F_+) = \frac{1}{Z(\mu)} \frac{F!}{F_+!(F - F_+)!} \exp \left[\frac{\mu}{2} ((2F_+ - F)^2 - F) \right]. \quad (3)$$

This gives the (total) probability attached to all cultural vectors with F_+ out of F traits marked as ‘+’ or ‘+1’, where μ is the parameter controlling the overall level of coupling between features. Moreover, $Z(\mu)$ is a normalization factor, namely the partition function in equation (D8). Note that $\sum_{F_+=0}^F p(\mu, F, F_+) = 1.0$, since the expression combines the probability of different possible configurations with the same F_+ , which, due to symmetry reasons are equally likely. There are $F!/(F_+!(F - F_+)!)$ such configurations (the physical ‘density of states’, where a ‘state’ would correspond to a possible configuration, rather than to a ‘cultural state’ used in the nomenclature of this study) for each F_+ .

The model is mathematically equivalent to the Ising model of magnetism on a fully connected lattice [27], described in the canonical ensemble, with the parameter μ replacing the ratio between spin-spin coupling and temperature, which controls for the overall level of alignment between spins. This parallel does not come as a surprise: for any statistical physics ensemble defined by the averages of certain, externally controlled/measured (physical) quantities, the mathematical derivation can be formulated in terms of maximum-entropy inference [26], which ultimately provides a statistical, information-theoretic justification of minimum-bias as a replacement for assumptions like ‘ergodicity’. Due to this parallel, the nomenclature related to spins is sometimes used instead of that related to cultural features.

Starting from equation (3), one can derive the expression for the correlation between any two features:

$$C(\mu, F) = \frac{1}{Z(\mu)} \sum_{F_+=0}^F \frac{(F-2)! ((2F_+ - F)^2 - F)}{F_+!(F - F_+)!} \exp \left[\frac{\mu}{2} ((2F_+ - F)^2 - F) \right], \quad (4)$$

based on the entire statistical ensemble. The details of this derivations are also given in appendix D.

In equations (3) and (4), the coupling parameter is positive: $\mu \in [0, \infty)$. Physically, this corresponds to ferromagnetism, meaning that alignment between spins is favored, a tendency which is enhanced with increasing μ . Using equation (3), one can check that, for vanishing coupling $\mu = 0$, the probability of choosing a configuration with a given F_+ is directly proportional to the number of such configurations, which is specified by the binomial coefficient preceding the exponential. As μ is increased, more emphasis is given to configurations with unequal numbers of −1 and +1 traits, at the expense of configurations that are more balanced. Using

equation (4), one can also check that the correlation $C(\mu, F)$ increases with increasing coupling μ , as expected, and that $C(0.0, F) = 0.0$ for any F .

3.2. The second scenario: groups

This section explains the S2G model, in the context of generating (sets of) cultural vectors in a stochastic way. This probabilistic model enforces an organization of cultural vectors in terms of two, equally sized groups, with high similarities within groups and low similarities between groups. The model defines a probability distribution p taking as support the same set of possible cultural vectors as in section 3.1: the cultural space defined by F binary features, with 2^F configuration. One of the groups is ‘centered’ around the configuration with a -1 trait with respect to each feature, while the other group is centered around the opposite configuration, having a $+1$ trait with respect to each feature. The model is designed such that all features contribute equally to the group structure. As a consequence, this induces a certain level of correlation for all pairs of cultural features. The strength of these correlations is controlled by the same free parameter that controls the strength of the group structure.

According to the S2G model, every cultural vector that is generated is first randomly assigned to one of the two groups, with equal probabilities. These two groups are denoted as the ‘ -1 ’ group and the ‘ $+1$ ’ group. Then, at the level of every feature, the trait is randomly and independently chosen from the two possibilities, but with different probabilities: the trait with the same sign as the group is chosen with probability $1 - 2\nu$, while the trait with the opposite sign is chosen with probability 2ν . Here, $\nu \in [0, 0.25]$ is the free model parameter controlling the strength of the group structure: lower ν values imply stronger group structure and stronger correlations between features, as made more explicit by equation (6) below. From this procedure, it follows that, at the level of every feature, each generated trait falls under one of the following situations:

- with probability $0.5 - \nu$, it is attached to a vector belonging to group -1 and has a value of -1 ;
- with probability ν , it is attached to a vector belonging to group -1 and has a value of $+1$;
- with probability $0.5 - \nu$, it is attached to a vector belonging to group $+1$ and has a value of $+1$;
- with probability ν , it is attached to a vector belonging to group $+1$ and has a value of -1 .

Note that the probabilities of the four cases add up to 1.0, that the combined probability of either value is 0.5 and that the probability of either group is also 0.5.

For this model, the probability that a generated configuration has F_+ traits $+1$ is:

$$p(\nu, F, F_+) = \frac{1}{2} \frac{F!}{F_+!(F - F_+)!} (2\nu)^{F_+} (1 - 2\nu)^{F - F_+} [(2\nu)^{F - 2F_+} + (1 - 2\nu)^{F - 2F_+}], \quad (5)$$

while the correlation between any two features is:

$$C(\nu) = 1 - 8\nu + 16\nu^2. \quad (6)$$

The mathematical derivations of equations (5) and (6) are given in appendix E. Note that the correlation in equation (6) behaves as expected, namely: $C(0.0) = 1$ (when the two groups are maximally dissimilar the correlation is maximal) and $C(0.25) = 0.0$ (when the two groups are indistinguishable the correlation is zero). Finally, equation (6) can be written in the form of a quadratic equation, whose solution reads:

$$\nu(C) = \frac{1 - \sqrt{C}}{4}, \quad (7)$$

after having taken into account that $\nu \in [0, 0.25]$. Note that the alternative, $1 + \sqrt{C}$ solution given by the quadratic formula would be valid for the $\nu \in [0.25, 0.5]$ interval, which is not used here, since it is entirely equivalent (up to an inversion) with the $\nu \in [0, 0.25]$ interval, while being relevant only if group -1 is allowed to be biased toward $+1$ traits instead of toward -1 traits, and viceversa, which is not the case here.

3.3. Mathematical comparison of the two scenarios

This section deals with the comparison between the FCI and the S2G models, in terms of properties that can be extracted directly from the equations in section 3.1 and section 3.2, without the need of randomly sampling from the two statistical ensembles. Specifically, we focus on the behavior of the feature–feature correlation (figure 3), the shape of the probability distribution (figure 4), the correlation-based matching between FCI and S2G (table 2), and the symmetry breaking phase transition (figure 5) associated to each model.

Figure 3 shows the behavior of the correlation between any two cultural features for the two models. First, figure 3(a) shows how the correlation entailed by the FCI model depends on the model parameter μ controlling the pairwise couplings between features, based on equation (4). Different curves correspond to different values of F . Note that the correlation increases from $C = 0.0$ to $C = 1.0$ as the coupling μ is increased, but it also increases as the number of features F is increased. Second, figure 3(b) shows how the correlation entailed by

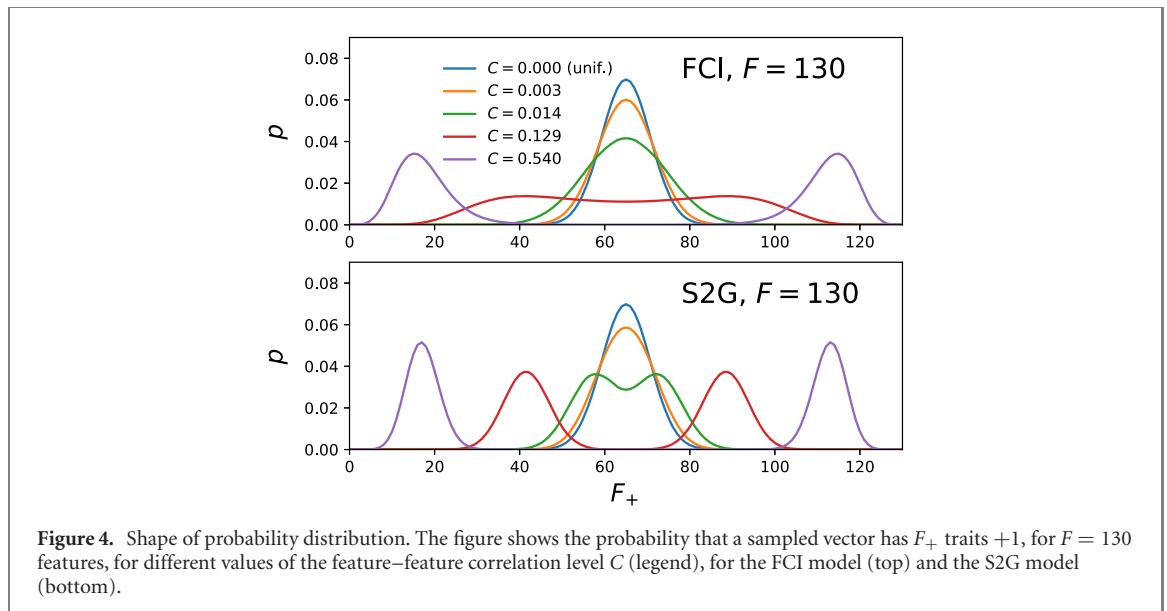
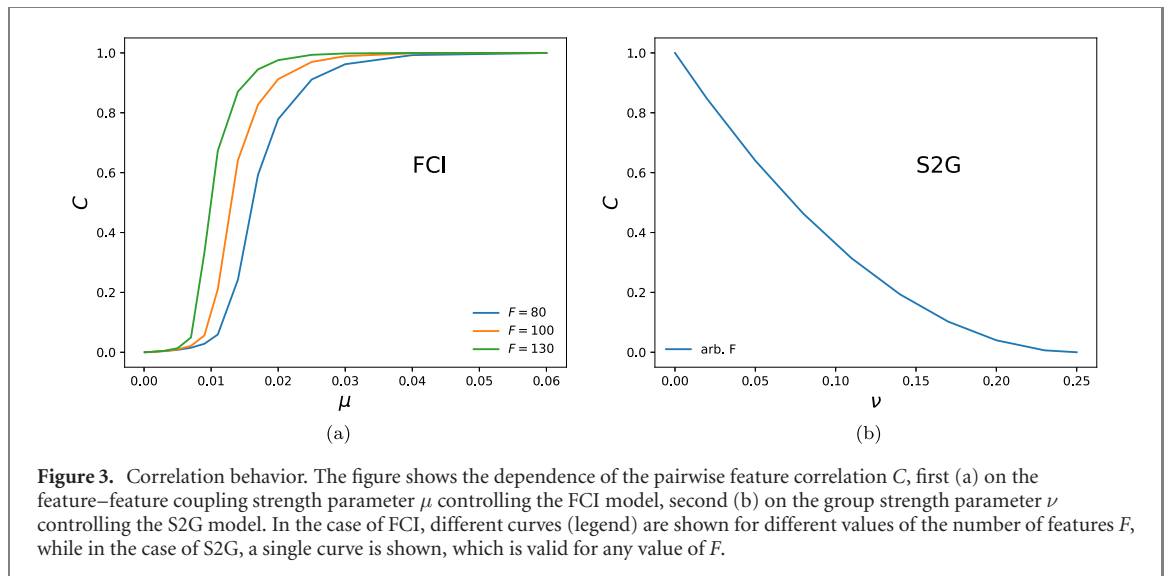


Table 2. Parameter mapping. The table shows the correspondence between the correlation values C shown in figure 4, the associated μ values used for generating the FCI probability curves and the associated ν values used for generating the S2G probability curves. This correspondence is valid when $F = 130$ features are used for the FCI and S2G models.

Correlation level C	FCI parameter μ	S2G parameter ν
0.000	0.000	0.250
0.003	0.002	0.237
0.014	0.005	0.221
0.129	0.008	0.160
0.540	0.010	0.066

the S2G model depends on the model parameter ν controlling the group strength, based on equation (6). Here, the correlation decreases from $C = 1.0$ to $C = 0.0$ as ν is increased, which is consistent with the fact that, by construction, lower values of ν correspond to a stronger group structure. Note that the $C(\nu)$ behavior is independent of F , which is obvious from equation (6).

All the following comparisons are based on a matching of the two models in terms of the correlation level C . Specific values of μ are chosen, based on which the correlation level entailed by the FCI model $C(\mu, F)$ is computed via equation (4), for a given F . Then, the corresponding $\nu(C)$ of S2G entailing the same correlation is calculated based on equation (7). This creates a correspondence between parameter μ of FCI and parameter

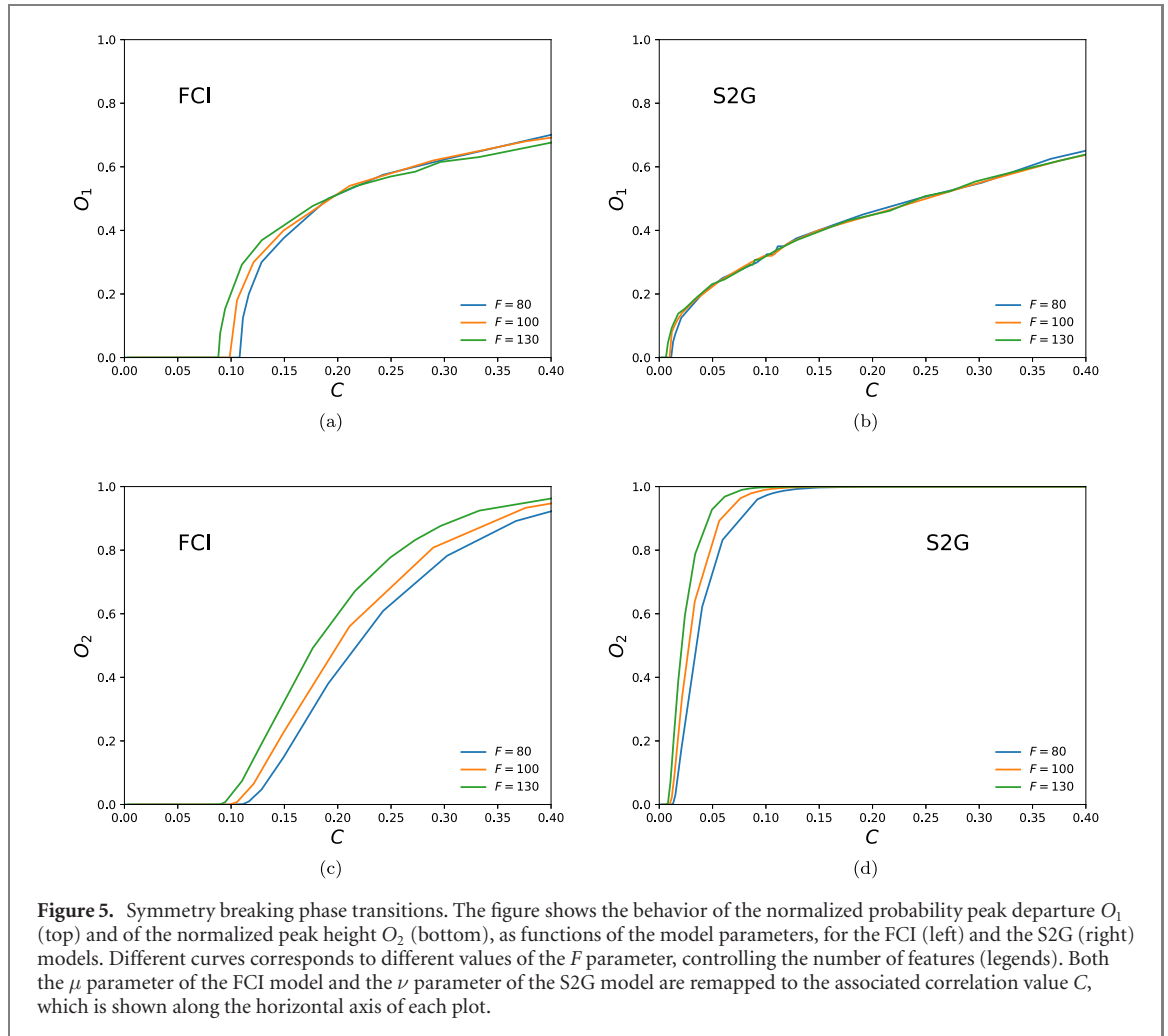


Figure 5. Symmetry breaking phase transitions. The figure shows the behavior of the normalized probability peak departure O_1 (top) and of the normalized peak height O_2 (bottom), as functions of the model parameters, for the FCI (left) and the S2G (right) models. Different curves corresponds to different values of the F parameter, controlling the number of features (legends). Both the μ parameter of the FCI model and the ν parameter of the S2G model are remapped to the associated correlation value C , which is shown along the horizontal axis of each plot.

ν of S2G by means of the correlation C . Since C is a number extracted from the full statistical ensemble under a specific parametrization, it can be regarded as a model parameter, namely as a replacement or remapping of μ (in the case of FCI) and of ν (in the case of S2G), which allows for a side-by-side comparison of the two models in terms of other quantities.

This μ -to- C -to- ν mapping is first exploited by figure 4, which shows the probability distributions associated to the FCI and S2G models, as described by equations (3) and (5) respectively. In either case, the distribution is shown for the same values of the correlation C that are listed by the legend at the top. These C values correspond to the values of the μ and ν parameters that are listed in table 2. The calculations are based on a value of $F = 130$, which is comparable to the F values associated to the empirical cultural states used in section 2 and section 5.

Note that, in the limit of vanishing correlation C , the distributions of both models converge to the uniform probability distribution, which assigns to every value of F_+ a probability that is equal to the fraction of possible configurations with that many ‘+’ traits. This uniform distribution is characterized by the existence of one maximum at the center of the F_+ axis. As the correlation C increases, the shape of the distribution becomes wider, with two equal maxima arising on either side of the F_+ axis, whose separation also increases with increasing C . Thus, each of the two models exhibits a symmetry breaking phase transition—a phase transition is a sharp boundary between two, qualitatively different regimes, occurring for a certain value of a model parameter, which becomes better defined and more obvious with increasing system size; the symmetry breaking aspect pertains to a certain type of change in the probability distribution when switching between the two regimes, which here is realized as a change between unimodality and bimodality; phase transitions are very important for statistical physics [28, 29] and complexity science [30].

Interestingly, a close inspection of figure 4 reveals that symmetry breaking occurs later (higher values of C) for the FCI model than for the S2G model, meaning that there is a non-vanishing C interval for which FCI exhibits unimodal behavior, while S2G exhibits bimodal behavior; interval which contains the $C = 0.014$ value. This C interval is of crucial interest for this study, since it corresponds to the correlation regime for which the symmetric group structure built into the S2G model is visible in the shape of the probability distribution,

while the feature–feature coupling built into the FCI model is not strong enough to induce a qualitatively similar shape. Still, even for C values that are high enough for the FCI distribution to also show maxima, the exact shapes of the two distributions are also different, with the S2G maxima being stronger than the FCI ones (visible for $C = 0.129$ and $C = 0.540$). This is a visual confirmation that the two statistical ensembles are indeed different and that the S2G ensemble has a smaller Shannon entropy than the FCI ensemble, for any, non-vanishing value of C , thus being more biased, more constrained and encoding more structure, which should manifest itself at the level of higher-order correlations (involving more than two spins/features).

A more complete picture of the phase transitions exhibited by the two models is provided by figure 5. This shows the dependence of two mathematical properties of the probability distributions in figure 4 on the model parameters. The first property, denoted here by $O_1(\gamma, F) \in [0, 1]$, is a normalized departure of either probability peak from the center of the (horizontal) F_+ axis. The second property, denoted here by $O_2(\gamma, F) \in [0, 1]$, is a normalized height of either probability peak compared to the probability at the center of the (horizontal) F_+ axis. Note that γ is a placeholder for either the μ parameter or the ν parameter, depending, respectively, on whether the FCI or the S2G model is used. Both quantities are zero when symmetry breaking is not present and are positive when symmetry breaking is present, giving higher values for better defined probability peaks. They can thus be used as ‘order parameters’ characterizing the phase transition, although they are evaluated in an *a priori* way, based on the expression of the probability distribution, rather than based on configurations sampled from the associated ensemble. Mathematically, the first quantity is defined as:

$$O_1(\gamma, F) = \frac{[0.5F] - F_+^*(\gamma, F)}{[0.5F]}, \quad (8)$$

while the second quantity is defined as:

$$O_2(\gamma, F) = \frac{p^*(\gamma, F) - p(\gamma, F, [0.5F])}{p^*(\gamma, F)}, \quad (9)$$

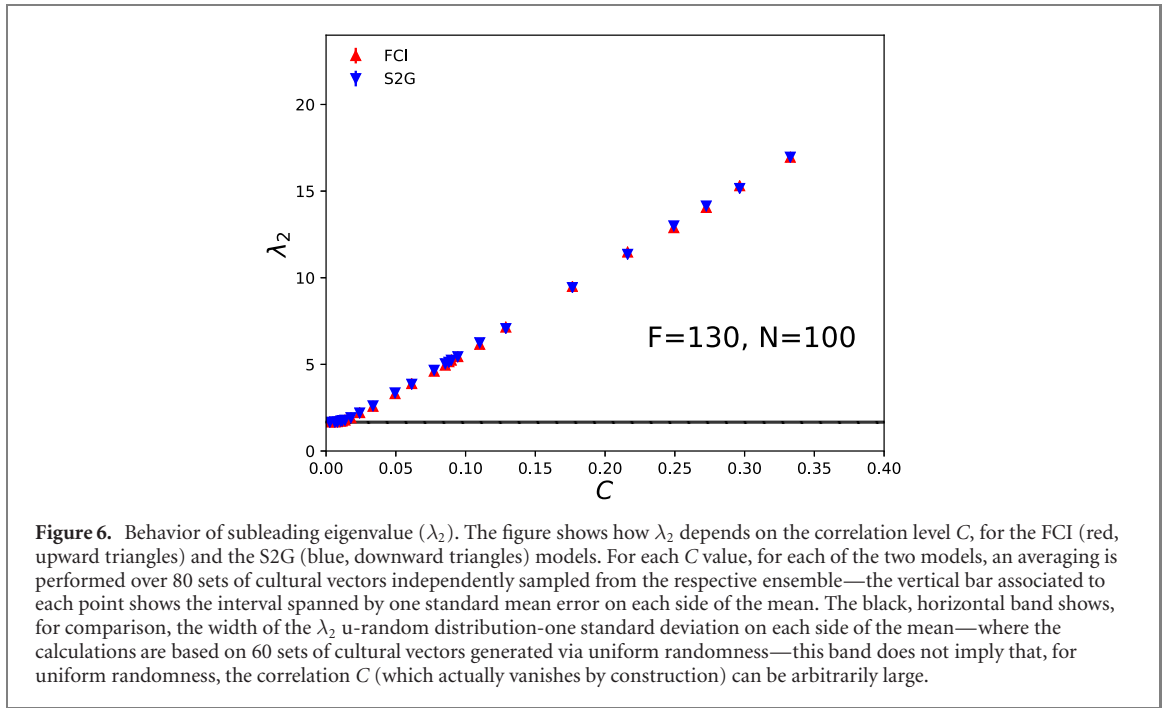
where the square brackets stand for the ‘integer part’ operation. Moreover, $F_+^*(\gamma, F)$ is the (integer) position along the F_+ axis of the first (lower- F_+) peak and $p^*(\gamma, F)$ is the height of this peak. At the same time, $p(\gamma, F, [0.5F])$ is evaluated according to either equations (3) or (5), depending on whether the quantity is evaluated for the FCI model (γ is replaced by μ) or for the S2G model (γ is replaced by ν). The value of $F_+^*(\gamma, F)$ is extracted by iteratively exploring the lower half of the F_+ axis, while evaluating $p(\gamma, F, F_+)$ according to either equations (3) or (5). On the other hand, $p^*(\gamma, F)$ is essentially an abbreviation for $p(\gamma, F, F_+^*(\gamma, F))$.

The four panels of figure 5 show the behavior of O_1 for the FCI model (figure 5(a)), of O_1 for the S2G model (figure 5(b)), of O_2 for the FCI model (figure 5(c)) and of O_2 for the S2G model (figure 5(d)). The dependence of either quantity on the μ parameter (for FCI) and on the ν parameter (for S2G) is translated in terms of the corresponding correlation value C , via equations (4) and (6) respectively. Note that the two quantities agree in terms of the correlation value for which the transition occurs, for both the FCI (figures 5(a) vs 5(c)) and the S2G (figures 5(b) vs 5(d)), for any number of features F . It is clear that the transition point comes closer to $C = 0.0$ with increasing F for both models. Finally, figure 5 shows that, independently of F , the transition point of S2G is located at lower values of C than that of FCI.

4. Discriminating between the two interpretations

This section investigates, from a spectral analysis and random matrix perspective, quantities that may differentiate between the two structural scenarios introduced in section 3: feature–feature redundancies vs individual grouping structure. To this end, sets of cultural vectors are numerically sampled from the two ensembles and similarity matrices are computed, based on equation (1). Since both the FCI and S2G ensembles are such that the (marginal) feature-level probability distributions are uniform, restricted randomness (see section 2) is equivalent to uniform randomness (see appendix A) as a null model (at least if the number of cultural vectors N is reasonably high) with respect to which structure is to be evaluated. Thus, for simplicity, uniform randomness (u-random) is used as a null model in this section. All comparisons made here make use of matching the feature–feature coupling parameter μ of FCI and the group strength parameter ν of S2G in terms of the correlation level C , as described in section 3.3. Moreover, the number of features and the number of cultural vectors are $F = 130$ and $N = 100$ for all the FCI, S2G and u-random cultural states generated and used for the figures of this section.

The most obvious quantity that could conceivably discriminate between the FCI and the S2G models is the subleading eigenvalue λ_2 , or the extent to which this goes above the uncertainty range predicted by uniform randomness. Figure 6 shows the dependence of λ_2 on the correlation level C for FCI (red) and S2G (blue), while the horizontal black band shows the u-random uncertainty range (1 standard deviation on each side of



the mean), as a compact replacement of the distribution in figure 17(a)—as mentioned in the figure caption, this band is not meant to give any information about the correlation level of the u-random null model, nor about realized correlations based on specific sets of vectors sampled from the ensemble. Surprisingly, λ_2 does not distinguish between the FCI and the S2G models, for any given correlation level C , since λ_2 values are identical, up to statistical errors arising from finite sampling. At the same time, λ_2 (for both models) does depart significantly from the null model expectations. This demonstrates that groups and redundancies are equally plausible explanations for empirical structural modes, such as those identified in section 2 (at least as long as the simplistic setting behind the FCI and S2G models is reasonably representative).

In the light of section 3 and appendix C, figure 6 also implies that the subleading eigenmodes of matrices produced via FCI are associated, on average, to the same self-similarity as those of matrices produced via S2G, for a given correlation level. This appears counter-intuitive, since the low- C presence of symmetry breaking for S2G makes it much easier to identify two, well separated groups, one for each side of the F_+ axis of figure 4. However, a closer inspection of the probability distributions in figure 4 reveals that FCI is more likely to produce, even in the absence of symmetry breaking, cultural vectors that are at one extreme or the other (almost fully populated with $+1$ traits or with -1 traits). These extreme configurations are much more representative, or ‘central’, for the configurations that are possible on the respective side of the F_+ axis, thus compensating for the softer separation at $F_+ = 0.5F$. Also note that the values of C used in figure 6 are the same for FCI and S2G and the same as those used in figures 7 and 8 described below. For each FCI and S2G point in any of these plots, explicit averaging over the sampled sets of cultural vectors is only performed with respect to the quantity associated to the vertical axes. For the correlation level C , associated to the horizontal axes, we simply use the analytically-computed, ensemble-level value, for the given parametrization of the model (equations (4) and (6)).

Appendix F shows, in a manner similar to figure 6, the behavior of the largest and third largest eigenvalues— λ_1 and λ_3 respectively—for the FCI and S2G models, in comparison to the u-random null model. The analysis there makes it clear that the λ_1 and λ_3 are both compatible with the null hypothesis. Thus, all or most of the structural information of cultural states generated from either the FCI or the S2G model is captured by the (λ_2, v_2) eigenpair. Since λ_2 cannot discriminate between the two scenarios, this means that all or most discriminating power is encoded in the associated eigenvector v_2 , which is the focus of the rest of this section.

Based on section 3.3 and in particular on figure 4, one can say that, for the interesting correlation interval where FCI does not exhibit symmetry breaking while S2G does, configurations that are on one side of the F_+ axis and are generated with S2G exhibit relatively equal fractions of traits of a certain sign, compared to those that are generated with FCI. The S2G configurations should thus also display relatively equal contributions to the structural mode (λ_2, v_2) , so the associated v_2 entries should be much more similar for S2G than for FCI. Given the symmetric nature of both models, it follows that the absolute values of all the v_2 entries should be much closer to each other for S2G cultural states than for FCI ones—in either case, the entries associated to

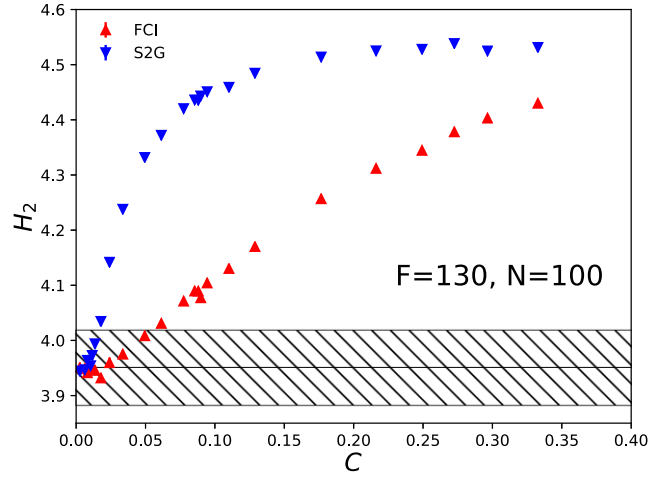


Figure 7. Behavior of uniformity H_2 associated to subleading eigenvalue. The figure shows how H_2 depends on the correlation level C , for the FCI (red) and the S2G (blue) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural vectors independently sampled from the respective ensemble—the vertical bar associated to each point shows the interval spanned by one standard mean error on each side of the mean. The black, horizontal band shows, for comparison, the width of the H_2 u-random distribution—one standard deviation on each side of the mean—where the calculations are based on 60 sets of cultural vectors generated via uniform randomness—this band does not imply that, for uniform randomness, the correlation C (which actually vanishes by construction) can be arbitrarily large.

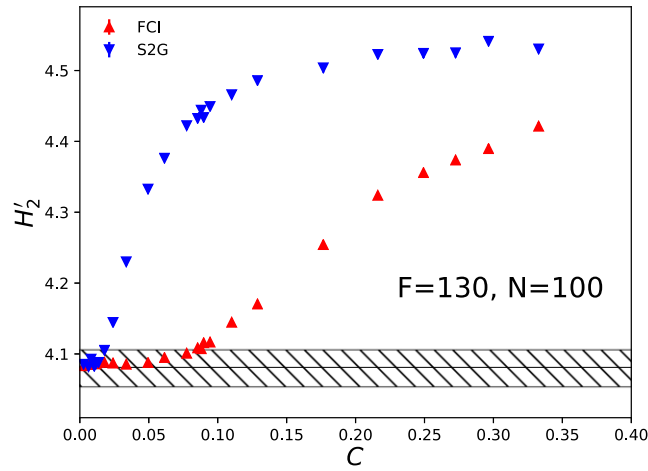


Figure 8. Behavior of subleading uniformity H'_2 . The figure shows how H'_2 depends on the correlation level C , for the FCI (red) and the S2G (blue) models. For each C value, for each of the two models, an averaging is performed over 80 sets of cultural vectors independently sampled from the respective ensemble—the vertical bar associated to each point shows the interval spanned by one standard mean error on each side of the mean. The black, horizontal band shows, for comparison, the width of the H'_2 u-random distribution—one standard deviation on each side of the mean—where the calculations are based on 60 sets of cultural vectors generated via uniform randomness—this band does not imply that, for uniform randomness, the correlation C (which actually vanishes by construction) can be arbitrarily large.

different sides of the F_+ axis would (typically) have different signs. This reasoning suggests that the difference between FCI and S2G would be captured by a quantity that evaluates the overall ‘uniformity’ of the v_2 eigenvector, based on the absolute values of its entries. Since these are normalized via $\sum_{i=1}^N |v_l^i|^2 = 1$ for any eigenvector v_l , the Shannon entropy is a natural quantity for evaluating the uniformity. This leads to the definition of the ‘eigenvector entropy’ H_l associated to the l th highest eigenvalue λ_l , as a measure of uniformity:

$$H_l = -\sum_{i=1}^N |v_l^i|^2 \log |v_l^i|^2, \quad (10)$$

where v_l^i is the i th entry of the eigenvector associated to λ_l —note that this quantity was also used in reference [22], which cites reference [31]; a variant of this quantity, involving a different type of normalization, is also used in reference [32], in relation to the global mode of empirical correlation matrices.

Figure 7 shows the behavior of the eigenvector entropy H_2 associated to the second highest eigenvalue λ_2 , in a format very similar to that of figure 6. This confirms that H_2 discriminates well between the two models, with

S2G showing clearly higher H_2 values than FCI as long as the correlation level does not come arbitrarily close to $C = 0.0$. Moreover, comparing the two profiles with the u-random one- σ band reveals that S2G becomes incompatible with the null-hypothesis for much lower correlation values than FCI. However, for either model, the $H_2(C)$ curve does not show the sudden increase that one would expect based on the phase transitions described in section 3.3, in the manner they are exhibited by the $O_1(C)$ and $O_2(C)$ curves in figure 5.

The smoothness of the $H_2(C)$ curves is related to the fact that, for the low- C regime, where λ_2 is highly compatible with the null hypothesis, H_2 is typically not the second highest eigenvector entropy, although it is associated to the second highest eigenvalue. This suggests a definition of H'_l as the l th highest eigenvector entropy, independently of the associated eigenvalue. Figure 8 is a modification of figure 7, with H'_2 used as a replacement for H_2 for the vertical axis, affecting all the FCI, S2G and u-random calculations. Note that, unlike in figure 7, the sudden changes in figure 5 are now reflected in figure 8. Moreover, the transition points at $F = 130$ in figure 5 seem to be well reproduced in figure 8, while the FCI and S2G shapes of the $H'_2(C)$ curves are quite similar to those of $O_2(C)$, which are related to the height of the probability distribution peaks. Finally, for higher C values, each $H'_2(C)$ curve in figure 8 is almost identical to the associated $H_2(C)$ in figure 7, so strong structure makes it very likely that the eigenvector of the second highest eigenvalue has the second highest entropy, and H'_2 is effectively equivalent to H_2 .

The considerations above strongly suggest that a significant departure of the eigenvector entropy from the null model expectation is a good indication that the eigenvector encodes information about a group or a grouping tendency: this departure is present for S2G but absent for FCI for the interesting C interval where S2G exhibits symmetry breaking and FCI does not. As expected, the criterion breaks down when C is higher than the critical FCI value, beyond which both models exhibit symmetry breaking and the associated distributions become increasingly similar with increasing C . This suggests that, in an empirical setting, the criterion is useful as long as the respective eigenmode is not too strong (in terms of eigenvalue and associated correlation level), although a generic, quantitative definition of ‘too strong’ is currently lacking and beyond the purpose of this study. As will be shown in section 6, there is another, more fundamental limitation of the uniformity criterion.

5. Revisiting the empirical data

The findings of section 4 highlight the importance of the eigenvector entropy, in addition to the eigenvalue, for deciding whether a structural mode qualifies as an authentic group mode or not. As a consequence, in this section, the two quantities are being used together for a second, more detailed inspection of the empirical data in section 2. The empirical similarity matrices are computed based on the same three sets of $N = 100$ cultural vectors used in section 2, constructed from EBM, GSS and JS data.

Figure 9 shows a scatter of the empirical eigenpairs of the EBM matrix, where the horizontal axis is associated to the eigenvalue λ , while the vertical axis is associated to the eigenvector entropy H . The global mode eigenpair is highlighted by the inset. The figure also shows the 1- σ uncertainty bands predicted by the r-random null model, for the subleading eigenvalue λ_2 , the leading eigenvalue λ_1 , the subleading entropy H'_2 and the leading entropy H'_1 .

The four structural modes identified via figure 2 are also visible in the main plot of figure 9, to the right of the vertical r-random band. Importantly, all their eigenvector entropies are below the horizontal r-random band, suggesting that neither of them qualifies as a group mode. Actually, all the bulk EBM eigenpairs are also below the r-random band, and thus compatible with the null hypothesis in terms of the uniformity of eigenvector entries.

The analysis in figure 9 is also applied to the other datasets and the results are presented in figure 10, with figure 10(a) showing the results for GSS data and figure 10(b) showing the results for JS data. In both cases, the results are similar to those of EBM data: the structural modes do not show a higher eigenvector uniformity than what is expected based on the null model, nor do any of the smaller- λ modes. In the light of section 3 and section 4, these results suggests that structural modes of empirical matrices of cultural similarity are not due to authentic group structure, but due to arbitrary semantic redundancies between the questions or items used for the dataset. However, such a conclusion would be premature, as it implicitly uses strong assumptions about cultural groups. This aspect is further explored in section 6.

Before proceeding to the section 6, another aspect is worth emphasizing. For all the empirical states, the leading eigenvector entropy is significantly smaller than what the null model predicts—although much higher than what is expected or realized for the structural modes and the random modes. This means that there is less equality among the contributions of different cultural vectors to the global mode than expected based on restricted randomness—although much more than for any of the structural modes. Thus, there is a significant level of heterogeneity in terms of individual identification with the system-level, ‘mainstream’ cultural tendency, independently of grouping heterogeneity potentially captured by the structural modes.

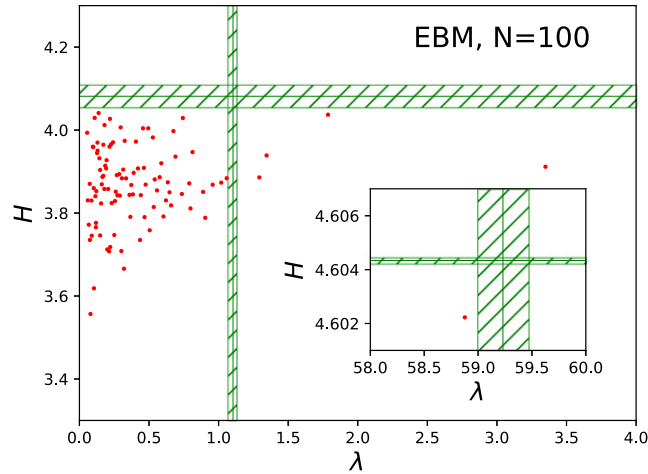


Figure 9. Eigenvalues and eigenvector entropies for empirical data. Every point corresponds to an empirical eigenpair, with the eigenvalue λ shown along the horizontal axis and the eigenvector entropy H shown along the vertical axis. The inset focuses on the leading eigenvalue, which also corresponds to the highest eigenvector entropy. The vertical bands in the main plot and in the inset show, respectively, the widths (one standard deviation on each side of the mean) of the subleading and leading eigenvalue distributions, based on restricted randomness. The horizontal bands in the main plot and in the inset show, respectively, the widths (one standard deviation on each side of the mean) of the second highest and highest eigenvector entropy distributions, based on restricted randomness. The vertical bands are not intended to provide any information about the eigenvector entropies associated to the respective eigenvalues, while the horizontal bands are not intended to provide any information about the eigenvalues associated to the respective eigenvector entropies. The figure is based on the same, EBM data with $N = 100$ cultural vectors used in figures 1, 2, 17 and 18.

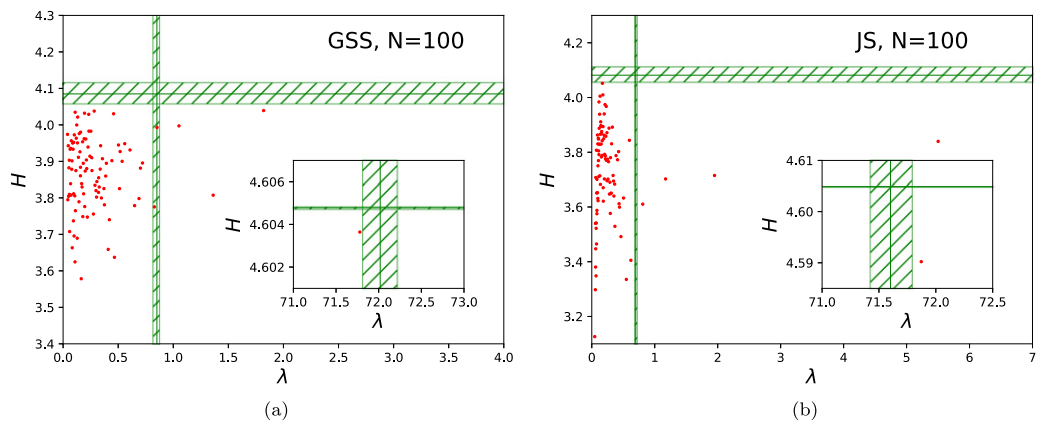


Figure 10. Eigenvalues and eigenvector entropies in other empirical datasets. (a) is based on the GSS data with $N = 100$ cultural vectors used for figures 2 and 19(a), while (b) is based on the JS data with $N = 100$ cultural vectors used in figures 2 and 19(b). Each plot makes use of the same type of eigenpair analysis as figure 9.

6. Internally non-uniform groups

This section elaborates on another interpretation of the above results: empirical structural modes (or at least some of them) are actually signatures of authentic, system specific groups that somehow do not exhibit significant eigenvector uniformity; this is an alternative to the interpretation resulting from section 5: empirical structural modes are due to arbitrary, instrument-specific feature–feature redundancies. The latter interpretation is based on the observation that the eigenvector uniformities of these modes are not higher than null model expectations, thus not satisfying the criterion emerging from section 3 and section 4. In turn, this criterion relies on comparisons between the S2G and FCI models and is potentially sensitive to assumptions about cultural groups that might be inherent in the S2G model, while not necessarily valid for real-world cultural groups.

In fact, the high subleading eigenvector uniformity characterizing S2G appears to be a direct consequence of the following property of the model: generated vectors associated to either group share a typical separation from what one may call the ‘center’ of the group—the hypothetical vector that would best represent that grouping tendency, which in this case would be either the full ‘−1’ or the full ‘+1’ vector. Most of the generated

vectors fall within a relatively narrow region around that typical separation. The separation bands associated to the two groups are obvious upon inspecting the bottom of figure 4, in the form of the two peaks of the probability distribution, present for a wide range of the correlation level C —in this representation, the two group centers correspond to the two extremes of the F_+ axis: $F_+ = 0$ and $F_+ = F$. As explained in section 4, such a well defined separation band implies that the eigenvector capturing the respective group (effectively capturing both groups, in the case of S2G, due to the highly symmetric setting) typically has a large number of entries with relatively similar (absolute) strengths and no entries of significantly higher strengths.

One can thus argue that, by construction, S2G gives rise to groups that are ‘internally uniform’, property which seems inadequate for real world cultural groups: it is hard to imagine a reason why individuals would be effectively forbidden from coming arbitrarily close to or from going arbitrarily far from the center of the group in cultural space. More concretely, if the group is a manifestation of an underlying ideology or way of thinking, there seems to be no reason why there should exist a preferred number of topics/items in terms of which individuals under the influence of that ideology would agree with its most representative opinion profile. This argument holds regardless of whether the grouping takes place around an ephemeral, contextual political movement, or around a more universal ‘rationality’ or ‘way of life’, like those postulated by reference [33]; the extent to which media-based or interpersonal influence is involved in sustaining or amplifying the grouping is also irrelevant for this argument.

Instead, it is very plausible that a real cultural group exhibits a high variability of the extent to which different individuals identify with it, so that one encounters non-vanishing numbers of individuals that are very central or very peripheral. Such ‘internally non-uniform’ groups would likely not exhibit statistically significant eigenvector uniformities, so the eigenvector entropy criterion developed and used above would fail to recognize them as authentic. In order to illustrate this scenario in a more quantitative manner, section 6.1 makes use of another toy model, called ‘mixed prototype generation’ (MPG), inherited from previous work [8], where it was shown to be capable of generating cultural states that reproduced other, important empirical properties. This model explicitly randomizes the strengths of the couplings between generated vectors and central group profiles, or ‘prototypes’, so that the distribution of vectors’ separations from these ‘prototypes’ is rather flat, without separation peaks/bands like those of S2G. The ‘mixing’ relates to the fact that every generated vector is a quasi-unique combination of all the prototypes, although typically dominated by either of them. While structurally different from both S2G and FCI, MPG can also be used in the binary-feature setting employed in section 3 and section 4 (although in a manner that is somewhat less elegant mathematically). This is exploited by section 6.1, which explicitly shows that cultural groups with strongly-significant eigenvalues but non-significant eigenvector uniformities may exist, thus providing an important, theoretical indication that empirical structural modes highlighted in section 2 and section 5 might still be signatures of authentic cultural groups that are internally non-uniform.

Another, more empirically-based indication of this possibility is provided by section 6.2. This takes advantage of the block-diagonal form of the feature–feature correlation matrix of one of the datasets used in this study, which allows for easy identification and elimination of blocks of obviously redundant features, while checking the robustness of the structural modes under this operation. It turns out that some of the structural modes retain their eigenvalue significance after eliminating all the obvious feature redundancies, suggesting that these robust modes are actually due to authentic but internally non-uniform cultural groups.

6.1. The mixed prototypes scenario

This section focuses on cultural states generated with the MPG procedure proposed in reference [8]. Besides its interesting social science foundation and demonstrated structural realism, MPG is capable of generating cultural states characterized by groups that are internally non-uniform, which is why it is employed here. After providing a review of the central assumptions and technical aspects behind MPG in the following few paragraphs, we present, via table 3 and figure 11, relevant random matrix results obtained for MPG cultural states. In parallel, figure 12 illustrates how MPG vectors are distributed with respect to the group prototypes, while providing further understanding of the essential structural differences between MPG, S2G and FCI.

MPG is designed [8] to generate cultural states that reflect the existence of a certain number K of underlying ideologies that are effectively recognizable, to different extents and in different ways, at the level of every individual in the population. These ideologies are formally represented by abstract cultural vectors, or prototypes, so that each prototype is the ‘ideal’ opinion profile of one ideology. Every concrete cultural vector is generated as a random, quasi-unique mixture of the K prototypes and a uniformly random, pure noise component. The associated $K + 1$ contributions are deliberately unequal, with the largest K being randomly attributed to the K prototypes, while the smallest is always associated to pure noise. Thus, for each vector, there is an arbitrary ordering of the prototypes in terms of the number of traits that are copied from each of them, while the number of traits generated from pure noise is by construction smaller or equal to that of the lowest-contributing

Table 3. Relevant estimates for mixed prototypes generation (MPG). The first column shows values of the MPG β parameter controlling the strength of prototype mixing. The second column shows values of the estimated feature–feature correlation level $\tilde{C}$, which is numerically computed by generating 20 000 MPG vectors for each β value. The last three columns show z-scores of the subleading eigenvalue $z(\lambda_2)$, the uniformity associated to the subleading eigenvalue $z(H_2)$ and the subleading uniformity $z(H'_2)$. The λ_2 , H_2 and H'_2 values are extracted from one MPG cultural state of $N = 100$ vectors generated for each β value, while the associated z-scores are computed with respect to the r-random null model, from which $n = 1000$ random matrices are sampled for each β value. All estimates are valid for $F = 130$ binary features. The highlighted values of β indicate direct correspondences with the plots in figure 11.

β	$\tilde{C}$	$z(\lambda_2)$	$z(H_2)$	$z(H'_2)$
0.20	0.00	0.41	−4.05	−0.33
0.40	0.04	20.92	−13.36	0.18
0.60	0.12	79.07	−6.24	−0.07
0.70	0.21	182.83	−4.50	−1.98
0.75	0.26	242.82	−1.44	−1.44
0.80	0.34	303.79	2.90	2.90
0.85	0.44	410.63	4.90	4.90
0.90	0.57	468.78	8.25	8.25

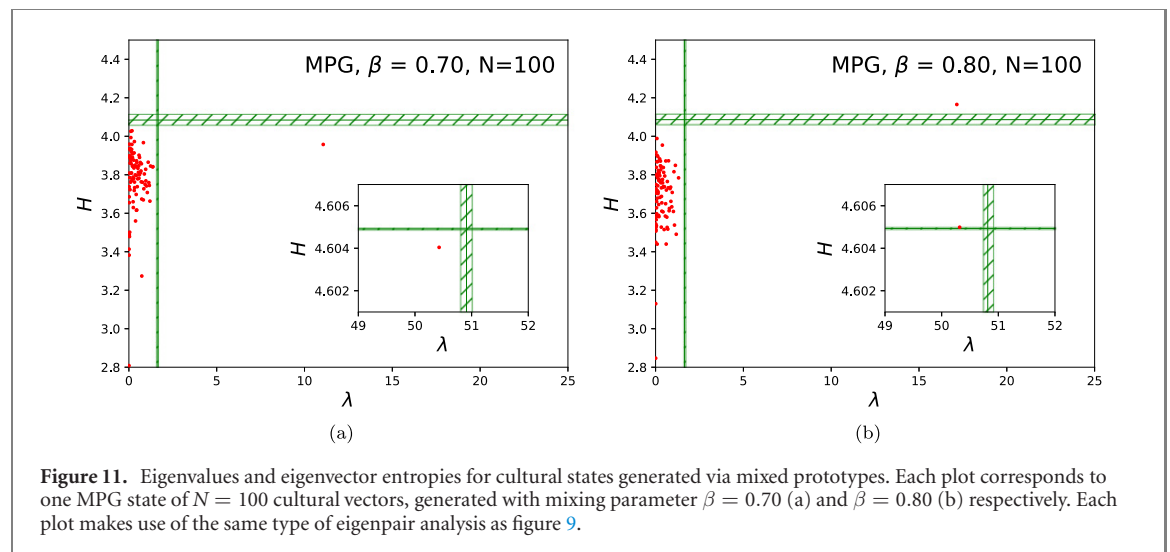


Figure 11. Eigenvalues and eigenvector entropies for cultural states generated via mixed prototypes. Each plot corresponds to one MPG state of $N = 100$ cultural vectors, generated with mixing parameter $\beta = 0.70$ (a) and $\beta = 0.80$ (b) respectively. Each plot makes use of the same type of eigenpair analysis as figure 9.

prototype. At the same time, each prototype is dominant (provides the largest contribution) for roughly $1/K$ of the generated vectors.

The mixing contributions are specified by randomly picking, for each generated vector, a set of $K + 1$ weights $\{w_1(\beta), \dots, w_{K+1}(\beta)\}$ satisfying $\sum_{i=1}^{K+1} w_i(\beta) = 1$, which are assigned to the K prototypes and to the pure noise component. The latter is bound to always receive the smallest weight, while the prototype that receives the highest weight is understood as the dominant prototype of that vector. For each vector, roughly $w_l(\beta)F$ features are randomly assigned to contribution l and the associated traits are generated accordingly, by either being copied from the respective prototype or being randomly generated, depending on the type of contribution that l stands for. Here, $\beta \in (0, 1)$ is a free model parameter controlling the overall/expected strength of the groups and mixing (stronger groups and weaker mixing for higher β), together with the shape of the joint probability distribution from which the weights are effectively sampled. This is essentially a β -dependent probability distribution taking as support the volume of the regular K -simplex spanned by a vector $\vec{w}$ taking the weights as its entries. MPG employs a pragmatic, computational sampling procedure that does not explicitly state nor use the associated joint distribution. While this procedure is explained in reference [8], it is worth providing here an intuition of the role that β plays: the joint distribution places most emphasis on the vertices of the simplex when $\beta \rightarrow 1$, while placing most emphasis to the center of the simplex when $\beta \rightarrow 0$. The former extreme corresponds to having one weight of almost 1.0, while the other are almost 0.0 (very strong groups and boundaries, very weak mixing). The latter extreme corresponds to having all weights almost equal to $1/(K + 1)$ (very weak groups and boundaries, very strong mixing).

In practice, for an intermediate β value, a generated MPG vector effectively falls under one of K possible types, or groups, depending on its dominant prototype. However, the flexibility of the weights associated to different prototypes make the boundaries between groups rather soft. Moreover, within each group, the vectors exhibit significant variability in terms of how close they actually are to the prototype, variability associated to the (β -dependent) marginal distribution of the largest weight. In turn, this gives rise to the internally non-uniform nature of MPG groups, containing vectors that are arbitrarily close to or far from the ‘core’.

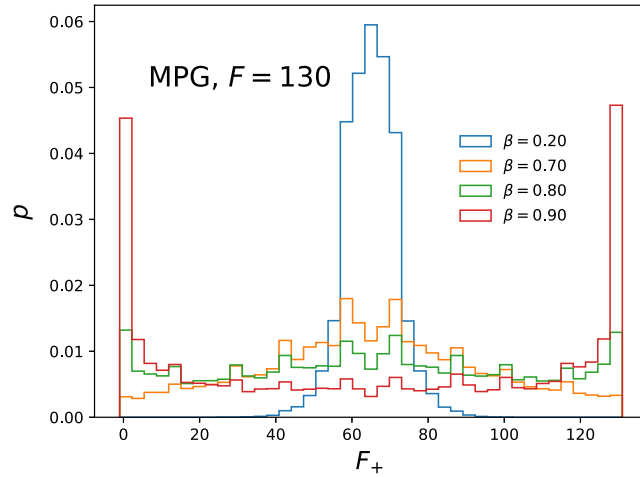


Figure 12. Approximate shape of probability distribution. The figure shows the probability that a sampled vector has F_+ traits $+1$, for MPG with two, maximally separated prototypes with $F_+ = 0$ and $F_+ = F$, for several values of the mixing parameter β (legend), for $F = 130$ binary features. The estimation is done based on $N = 20\,000$ numerically sampled vectors for each β value. Note that the histograms use 41 bins that cover a probability support with $F + 1 = 131$ discrete points, so each bin includes either 3 or 4 of these points.

The above description of MPG is conditional on cultural prototypes already being fully specified in terms of the traits they pick for every feature. In reference [8], the K prototypes themselves are also randomly generated during a preliminary step, according to a procedure that uses another parameter α controlling for the separation between prototypes. Here, instead, $K = 2$ prototypes are manually defined and kept fixed, in a manner that allows for direct comparisons with the S2G and FCI models introduced and studied in section 3 and section 4. Specifically, a cultural space consisting of $F = 130$ binary features is used, where the two (maximally-dissimilar) prototypes are filled entirely with ‘ -1 ’ and ‘ $+1$ ’ traits respectively (the same labeling convention as in section 3), so that a binary, symmetric group structure is induced. These prototypes coincide with the central/representative vectors of the two S2G groups, as well as with the unique spin configurations corresponding to the two extremes of the F_+ axis in figure 4. Although S2G also has a binary, symmetric group structure, the S2G probability of generating vectors identical with the extreme configurations effectively vanishes, as long as the group strength parameter ν is not too low (group strength and correlation level not too high). As suggested by the above explanations and shown by figure 12, this is not the case for MPG. Just like for FCI and S2G, the binary, symmetric setting used for MPG gives rise to the expectation that only one structural mode would be present, with an associated λ_2 eigenvalue increasing with increasing β . Table 3 and figure 11 confirm this expectation.

A multitude of values are selected for the β parameter (listed in table 3) and one cultural state is generated for each of these values. For each of these states, the enhanced eigenpair analysis from section 5 is carried out, showing the presence of only one structural mode, whose eigenvalue increases with β . The two plots in figure 11 illustrate this analysis for $\beta = 0.70$ and $\beta = 0.80$. For these values, the structural mode is, respectively, below and above the r -random expectation, in terms of eigenvector uniformity. Note that this expectation band is crossed for a very high (and extremely significant) subleading eigenvalue $\lambda_2 \approx 15$. For comparison, the FCI and S2G models exhibit similar crossings already when $\lambda_2 \approx 5$ and $\lambda_2 \approx 1$ respectively, as revealed when combining the information in figures 6 and 8. The comparison between MPG and S2G confirms that the binary group structure induced by the former is very different from that induced by the latter, as the structural mode needs to be almost 15 times stronger in order to exhibit significant eigenvector uniformity. The comparison between MPG and FCI confirms that the eigenvector uniformity criterion is entirely inappropriate for validating groups like those induced by MPG, as the structural mode needs to be roughly 3 times stronger in order to exhibit significant eigenvector uniformity—the criterion is much less likely to (correctly) identify structural modes based on mixed prototypes as authentic than (erroneously) identify structural modes based on feature–feature redundancies as authentic.

Such insights become even clearer when inspecting table 3, which shows how several relevant quantities depend on the MPG mixing parameter β . In particular, the second column shows a numerical estimate $\tilde{C}$ of the feature–feature correlation level, based on $N = 20\,000$ MPG vectors generated for each value of β . The value of $\tilde{C}$ is obtained by averaging over the $F/2 = 65$ Pearson correlation values corresponding to all pairs of consecutive features—although this estimator might be biased, it should be consistent (asymptotically approach the true value in the limit of $N \rightarrow \infty$); unlike FCI and S2G, MPG does not seem to allow for an

exact, analytic, full-ensemble formula for computing C . The following three columns show the z -scores of the subleading eigenvalue λ_2 , the associated eigenvector entropy H_2 and the subleading eigenvector entropy H'_2 . For each β value, λ_2 , H_2 and H'_2 are extracted from one MPG cultural state of $N = 100$ vectors (the MPG state used in figure 11, when $\beta = 0.70$ and $\beta = 0.80$), while the associated z -scores are computed with respect to the subleading eigenvalues (for λ_2) and subleading eigenvector entropies (for H_2 and H'_2) of $n = 1000$ r-random matrices suitable for the respective MPG state (the same r-random matrices on which the uncertainty bands of figure 11 are based, when $\beta = 0.70$ and $\beta = 0.80$).

One notices in table 3 the increase of the eigenvalue significance $z(\lambda_2)$ of the MPG structural mode with increasing β and increasing correlation $\tilde{C}$. Only for $\beta = 0.75$ does the eigenvector entropy of the structural mode qualify as the subleading eigenvector entropy, while still smaller than the r-random expectation, since $z(H_2) = z(H'_2) < 0$. When $\beta = 0.80$, the structural mode also becomes significant in terms of eigenvector entropy, with respect to the r-random expectation, as $z(H_2) = z(H'_2) > 2.0$, while this significance further increases for higher β values. A transition takes place somewhere between $\beta = 0.75$ and $\beta = 0.80$, so that for higher β the two, symmetric groups induced by MPG, captured by the λ_2 structural mode start exhibiting internal uniformity that becomes statistically recognizable via eigenvector entropy. Note that the correlation-level $\tilde{C} \approx 0.2$ corresponding to this transition is much higher than that associated to the S2G phase transition $C \approx 0.01$ (figure 5-right and figure 8-blue) and higher even than that associated to the FCI phase transition $C \approx 0.08$ (figure 5-left and figure 8-red).

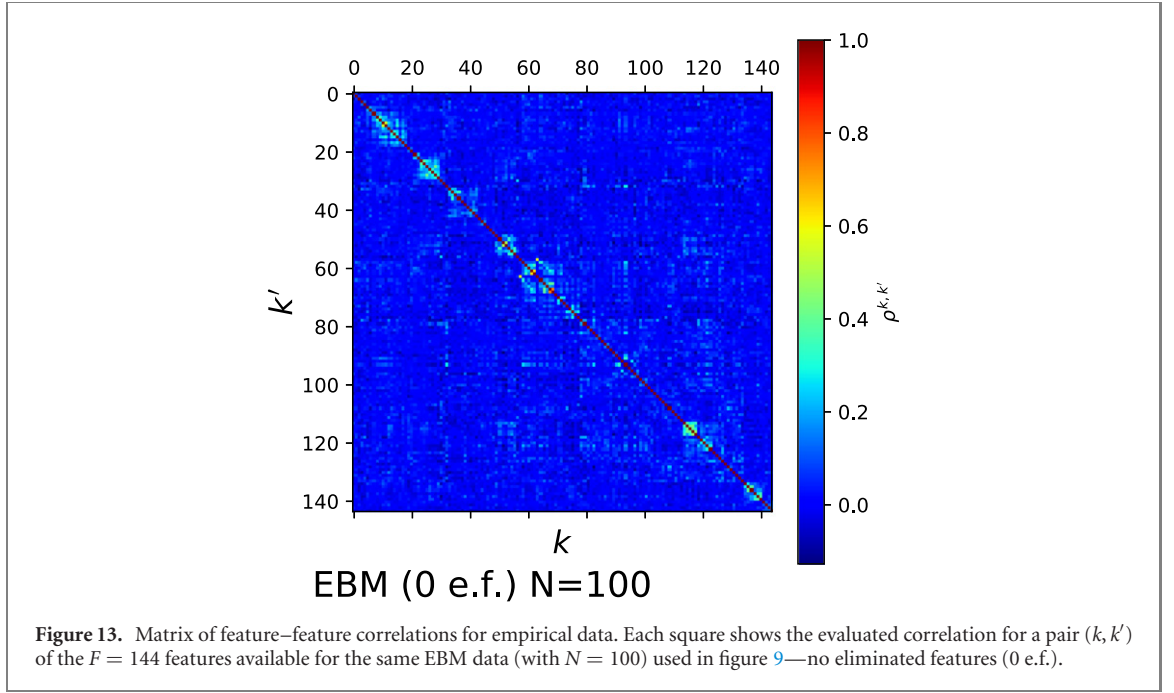
With these results in mind, it is instructive to see how MPG vectors are distributed in terms of their number F_+ of ‘+1’ traits. This is shown in figure 12, for several β values associated to MPG states described by table 3, two of which ($\beta = 0.70$ and $\beta = 0.80$) are also present in figure 11. As a reminder, each of the two extremes ($F_+ = 0$ and $F_+ = F$) of the horizontal axis is compatible with only one possible configuration of traits, which perfectly matches one of the two prototypes, so the location on the axis also determines the separation from the two prototypes: $\max(F_+, F - F_+)$ is effectively the similarity with the dominant prototype of a vector sitting at F_+ . As expected, the three higher β values, which exhibit a significant λ_2 (see $z(\lambda_2)$ column in table 3), also show F_+ distributions that are rather broad and flat, when compared to the S2G (and even to the FCI) distributions in figure 4. Although such comparisons are somewhat obscured by MPG distributions being shown for different correlation levels than S2G (and FCI) distributions (a correlation-based correspondence like that inherent in figure 4 is hard to achieve, since analytic calculations are not feasible for MPG), the trend is clear: MPG groups are internally much more flexible than those of S2G, in terms of the separation of the vectors from the group cores, which makes them exhibit low eigenvector uniformity. Compared to S2G and even to FCI, the MPG correlation level needs to be much higher in order to exhibit a bimodal-like F_+ distribution, in the form of an accumulation of vectors very close to the two prototypes, which is visible for $\beta = 0.80$ and obvious for $\beta = 0.90$. This accumulation is responsible for the significant eigenvector uniformity exhibited by states with $\beta = 0.80$ or higher.

On one hand, one might also notice the presence of two, small probability peaks close to the center of the F_+ axis in figure 12, for $\beta = 0.70$ and $\beta = 0.80$. This is likely a consequence of MPG being formulated in a somewhat arbitrary way (the lack of an explicit, analytic control of the joint weight distribution, the presence of the pure noise component), which is inherited from reference [8]—future research on internally non-uniform groups and/or mixing prototypes would likely benefit from a revised version of MPG. In any case, these peaks cannot drive up the uniformity, since the eigenvector entries of associated configurations are relatively weak: there are many other configurations further to the extremes of the axis, which are closer to the group cores and to each other. By contrast, the symmetry-breaking peaks of S2G and FCI (figure 4) do drive up the uniformity, since there is a vanishing number of configurations further to the extremes, even if the peaks themselves also arise quite close to the center. On the other hand, the smaller discontinuities in the shapes of the MPG distributions are due to fluctuations inherent in the numerical sampling on which the estimation is based and the histogram binning: $n = 20\,000$ sampled vectors are divided among 131 values of F_+ which are divided among 41 bins.

It is thus possible to construct groups that are internally strong but also non-uniform, which exhibit high and significant eigenvalues but low and non-significant eigenvector entropies. Such groups would not be recognized as authentic by the analysis applied in section 5, although they are more plausible as manifestations of real-world ideologies than internally-uniform groups. This is clearly illustrated by the mixed prototypes scenario, which builds on theoretical considerations from social science, while also exhibiting properties that are generically compatible with empirical data [8].

6.2. Redundant feature elimination

We have repeatedly emphasized during previous sections that empirical cultural states exhibit arbitrary feature–feature redundancies that have to do with how the underlying survey is designed rather than with system-specific properties. As shown in previous work [5–7], such redundancies are often visible at the level of the



feature–feature correlation matrix. Although correlations between features can be due either to groups or to redundancies (see section 3), for some datasets, subsets of obviously redundant features may be easily identified by inspecting the correlation matrix. This allows for eliminating these obvious redundancies and investigating the stability of structural modes under this operation. Stable structural modes are much more likely to be due to system-specific groups.

This section focuses on the EBM dataset, for which the redundancies between features are most obvious. We start by illustrating this with figure 13, showing the matrix of correlations between the $F = 144$ features, based on the $N = 100$ cultural state used above. Notice the block-diagonal form of the matrix, signaling the presence of multiple blocks of consecutive features that are highly correlated with each other. The features within each block actually correspond to survey questions that are concerned with different aspects of the same topic. For instance, within the block corresponding to features $k, k' \in [9, 19]$, the items measure attitudes with respect to 11 policy proposals that were part of the Maastricht Treaty, which aimed at enhancing integration within the European Union. Thus, these blocks are clearly due to survey-dependent redundancies between features.

It is worth mentioning at this point that the matrix entries in figure 13 are not computed based on the standard, Pearson correlation formula, since this is not appropriate when both nominal and ordinal variables are present. Instead, the values are based on a variant of Pearson correlation that uses the feature-level similarity values associated to different pairs of cultural vectors, rather than using actual traits attained by different vectors. Formally, for a pair of features (k, k') , this modified correlation reads:

$$\rho^{k,k'} = \frac{\sigma^{k,k'}}{\sqrt{\sigma^{k,k} \sigma^{k',k'}}}, \quad (11)$$

where the associated variance-covariance matrix is given by:

$$\sigma^{k,k'} = \langle s_{ij}^k s_{ij}^{k'} \rangle_{i,j \in 1,N}^{i < j} - \langle s_{ij}^k \rangle_{i,j \in 1,N}^{i < j} \langle s_{ij}^{k'} \rangle_{i,j \in 1,N}^{i < j}, \quad (12)$$

where s_{ij}^k is the similarity associated to a single feature k —its expression is visible upon eliminating the averaging over $1, \bar{F}$ in equation (1)—and in all cases the averaging is performed over all distinct pairs (i, j) of cultural vectors. This modified correlation is effectively identical to that used in references [5, 7], where it is formulated in terms of (feature-level) cultural distances d_{ij}^k instead of similarities s_{ij}^k , formulation which is mathematically equivalent, due to the simple, linear relationship ($d_{ij}^k = 1 - s_{ij}^k$) between distances and similarities. Because of the unconventional formulation, the correlation values in this section cannot be directly compared to those in section 3, section 4 and section 6.1, which rely on a conventional Pearson formulation, which is appropriate for a cultural space built entirely from binary features.

The next step is to eliminate features from the EBM dataset so that the redundancy blocks in figure 13 are no longer present. We describe here a deterministic procedure/algorithm that does this in a sequential way, so that a specified number of features are eliminated one by one. Let the (dynamic) set of features $\mathcal{F}$ initially

Table 4. Robustness of structural modes under redundant feature elimination. The first column shows the number of eliminated features (n.e.f.). Each of the following four pairs of columns corresponds to one of the structural modes in figure 9. The two columns of each pair show, respectively, the eigenvalue $z(\lambda_l)$ and the eigenvector entropy $z(H_l)$ associated to the l th empirical eigenvalue. The z -scores are computed with respect to the r -random null model, from which $n = 1000$ random matrices are sampled for each number of eliminated features. The calculations are based on the EBM data with $N = 100$ used in figure 9.

n.e.f.	$z(\lambda_2)$	$z(H_2)$	$z(\lambda_3)$	$z(H_3)$	$z(\lambda_4)$	$z(H_4)$	$z(\lambda_5)$	$z(H_5)$
0	70.38	-6.24	21.37	-1.68	7.54	-5.25	5.89	-7.17
10	62.48	-8.05	20.42	-1.45	5.49	-1.61	3.17	-7.09
20	55.08	-9.44	17.20	-1.34	3.31	-5.89	1.19	-7.64
30	47.35	-10.80	13.68	-3.06	2.65	-3.71	0.71	-6.61
40	40.86	-9.44	8.23	-3.47	0.48	-3.16	-0.07	-6.99
50	33.91	-10.40	6.13	-2.89	0.78	-0.40	-0.58	-6.74
60	30.00	-11.67	5.58	-3.66	0.50	-1.11	-2.14	-3.24

contain the integer labels of all features: $\mathcal{F} = \overline{1, F}$. At each step, the feature $k^* \in \mathcal{F}$ that is ‘most correlated’ is identified, according to the following criterion:

$$k^* = \operatorname{argmax}_{k \in \{k_1, k_2\}} \left[\max_{k' \in \mathcal{F} - \{k_1, k_2\}} \rho^{k, k'} \right], \quad (13)$$

where k_1 and k_2 are jointly defined by:

$$(k_1, k_2) = \operatorname{argmax}_{(k, k')} \rho^{k, k'}, \quad (14)$$

where $k, k' \in \mathcal{F}$ and $k < k'$. Feature k^* is then eliminated from the dynamical set: $\mathcal{F} := \mathcal{F} - \{k^*\}$. The procedure continues with the next step, unless the desired number of eliminated features (n.e.f.) has already been achieved. In a less formal language, at each step in the iteration, from the set of surviving features $\mathcal{F}$, one identifies the pair of distinct features (k_1, k_2) exhibiting the largest correlation—equation (14). Among these two features, one eliminates the one that exhibits the largest correlation with any other surviving feature different from k_1 and k_2 —equation (13). Note that one can think of other criteria for identifying ‘the most correlated’ feature k^* at each step in the algorithm, since each feature will generally exhibit a different correlation value with each of the other features in $\mathcal{F}$. The criterion described by equations (13) and (14) represents a pragmatical, greedy-type approach that we believe is suitable for the current analysis.

Table 4 illustrates the behavior of interesting EBM eigenmodes upon gradually increasing the number of features eliminated with the above procedure. The focus is on the eigenmodes associated to eigenvalues $\lambda_2, \lambda_3, \lambda_4, \lambda_5$, which are the EBM structural modes (see figure 9) when the n.e.f. (first column) is still zero (first line). The table shows the statistical significance (z -scores) of the eigenvalue and eigenvector entropy associated to these modes, computed with respect to the r -random null model. In terms of eigenvalue z -scores, one can see that λ_5 and λ_4 become compatible with r -randomness when n.e.f. reaches a value of 20 and 40 respectively, while λ_2 and λ_3 remain incompatible with restricted randomness even when n.e.f. reaches 60—note that the block diagonal form of the correlation matrix is no longer recognizable after eliminating 60 features, as shown by figure 14. In terms of eigenvector entropy, as expected, all four modes remain compatible with r -randomness, as indicated by the negative values of the associated z -scores. These results suggest that the two weakest EBM structural modes (λ_4, λ_5) are artifacts of feature redundancies, while the two strongest ones (λ_2, λ_3) are signatures of authentic grouping tendencies, although they (consistently) fail to exhibit any eigenvector uniformity.

At this point, it is essential that the statistical significance of the two, stronger structural modes is still clear after all the obvious feature redundancies are eliminated. It is much less important that their eigenvalue significance decreases to a certain extent under feature elimination, which is actually compatible with what one would expect from authentic structural modes. On one hand, this has to do with a decrease of the eigenvalues which is conceivably due to the fact that some features that are eliminated also store information about authentic cultural groups, despite exhibiting strong redundancies with other features (note that feature elimination is not carried out based on noisiness considerations). On the other hand, this has to do with an increase of the upper boundary of the random bulk—in this case, the r -random expectation for λ_2 —when reducing F while keeping N constant, which is compatible with naive extrapolations from the (much better understood) behavior of random correlation matrices. Both aspects seem jointly responsible for the effective decrease in ‘discrimination power’ encoded by the decrease of eigenvalue z -scores: the presence of both effects

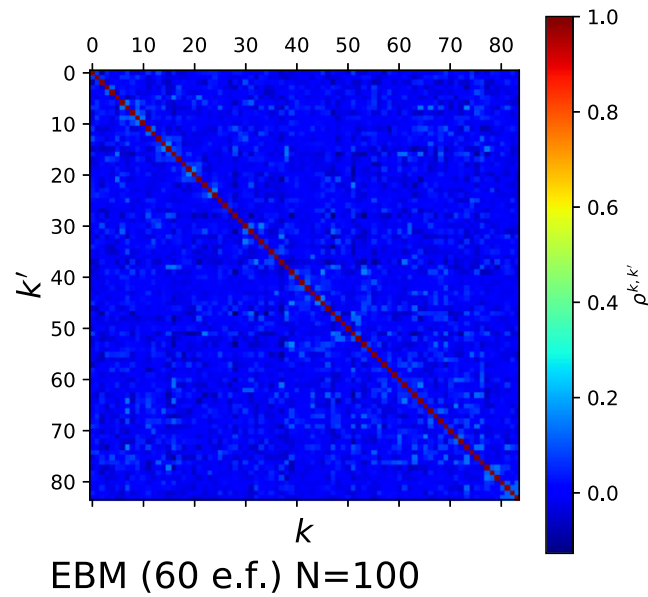


Figure 14. Matrix of feature–feature correlations for empirical data after extensive elimination of redundant features. Each square shows the evaluated correlation for a pair (k, k') of the $F - 60 = 84$ features available for the EBM data (with $N = 100$) used in figure 9, after the 60 most redundant features are eliminated (60 e.f.).

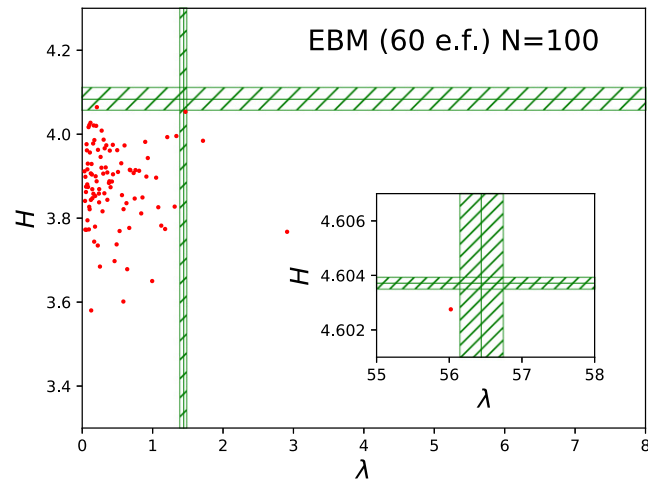


Figure 15. Empirical eigenvalues and eigenvector entropies after extensive elimination of redundant features. The figure illustrates the same type of eigenpair analysis as figure 9, for the same EBM data with $N = 100$ cultural vectors, after eliminating the 60 most redundant of the $F = 144$ features available initially (60 e.f.).

is confirmed by figure 15, which shows smaller empirical λ_2 and λ_3 and higher r-randomness λ_2 than figure 9 (before feature elimination).

It is important to emphasize that the feature-elimination procedure used in this section is mostly meaningful for the current context and purpose: using EBM data, it aids the idea that structural modes identified by random matrix theory can be authentic, even if they do not exhibit an eigenvector uniformity that is higher than the null model expectation. The procedure does not qualify as a general approach for validating structural modes, since many datasets affected by redundancies might not exhibit an obvious, block-diagonal structure of the feature–feature matrix. Moreover, if the survey designer has *a priori* intuition about and interest in the cultural groups that exist in the society that is being measured, the variables might actually exhibit (due to deliberate design or unconscious bias) a grouping into subsets that are associated to different cultural groups, which could induce a block-diagonal structure in the feature–feature matrix. However, the features in a block would then retain much valuable information about a system-specific group, so eliminating them would be counterproductive. In such a situation, the authentic-vs-artifactual question translates to a question of whether the latent construct behind a certain block is well aligned with a systemic grouping tendency or not. Although very interesting, such problems and questions are beyond the purpose of the current study.

7. Discussion

When it comes to the analysis of cultural states, this is the first study that employed a random matrix perspective, allowing for a separation of structurally irrelevant eigenmodes from the structurally relevant ones. When it comes to complex systems applications of random matrix theory, this is probably the first study that did not make use of an underlying time dependence: normally, aggregating over time provides the basis for evaluating, in most cases, correlations between the units of a system, or, in some cases, other association measures, like histogram-based mutual information [34] and Hamming-type similarity [22]—references [35, 36] provide recent examples where the latter is also used in time-dependent contexts, but independently of random matrix theory. Instead, the aggregation here involved a multitude of properties that are simultaneously measured for all units. Another novel aspect is the repeated numerical sampling of random matrices from the null model, which provided a grip on the uncertainty of the boundary between the noisy and the structural spectral regions, and thus a grip on the statistical significance of deviating eigenvalues. Another unique aspect is the extensive focus on the question of whether the structural modes identified in section 2 are just signatures of measurement-specific redundancies or, more interestingly, signatures of system-specific groups.

We started tackling this question with the help of two toy models, named ‘FCI’ (section 3.1) and ‘S2G’ (section 3.2), implementing the ‘redundancies’ and ‘groups’ scenario respectively. Each model has one parameter that controls is structural strength and, consequently, the level of correlations between features. We have shown that there is a non-vanishing correlation interval (section 3.3), with well defined boundaries (via the symmetry breaking phase transitions of S2G and FCI on the low and high correlation sides respectively), for which the two models induce very different probability distributions in cultural space (S2G exhibits bimodality, while FCI does not). In spite of these discrepancies, the two models were shown to exhibit exactly the same $\lambda_2(C)$ profile (section 4)—the same increase of the deviating eigenvalue with correlation—for all correlation values. This confirmed, in a surprisingly strong way that, at least in principle, deviating eigenvalues can be due to either redundancies or groups, in an equally plausible manner. At the same time, section 4 showed that complementary information about eigenvector uniformity is what distinguishes between FCI and S2G: for the interesting correlation interval, the subleading uniformity is compatible with null model expectations for FCI but significantly higher for S2G. Section 5 applies these theoretical insights to the empirical cultural states in section 2, showing that all non-leading eigenvector uniformities are compatible with null model expectations, including those associated to the previously identified structural modes. The first impression is that one can entirely reject the existence of cultural groups, along with the ‘prototypes’ hypothesis previously showing promising results [6, 8]. However, as explained in section 6, it is very plausible that real world groups are internally non-uniform, unlike those generated by S2G, and thus not recognizable via the eigenvector uniformity criterion. The internal non-uniformity property is strongly related to the ‘prototype mixing ingredient’, which had previously been essential for generating realistic cultural states [8]. This argument became more specific and quantitative in section 6.1, which made explicit use of the MPG model introduced in reference [8]. Then, section 6.2 presented a complementary argument in favor of the idea that some of the empirical structural modes correspond to authentic but internally non-uniform groups, by eliminating feature redundancies that are easily identifiable in the EBM dataset.

In relation to the MPG-based analysis in section 6.1, one could be tempted to interpret the absence of two, well defined peaks (like those of S2G) in the F_+ distribution (figure 12) as an absence of groups below the critical C value. Indeed, the results of section 3 and section 4 may seem to implicitly suggest a (re)definition of groups via peaks in the F_+ distribution. We argue that this interpretation is not appropriate. On one hand, even for low C values, the MPG distribution gives significantly more emphasis to the extremes of the F_+ axis than either S2G or FCI. On the other hand, F_+ values that are closer to these extremes carry much fewer possible configurations which are, on average, more similar to each other. This means that MPG generally induces a relatively high density in cultural space around these extremes (the prototypes), effectively providing ‘hard cores’ for its groups, which otherwise have relatively ‘soft external boundaries’. By contrast, S2G induces groups with ‘hard external boundaries’ and ‘hollow cores’, which are thus easily recognizable via eigenvector uniformity. Thus, MPG comes with a somewhat different meaning for ‘groups’ and ‘group structure’ than that implicit in S2G, which is conceptually more compatible with what one would expect from cultural groups in the real world, regardless of whether they are centered on universal rationalities or on more contextual, ephemeral ideological movements, as explained at the beginning of section 6.

One may also wonder, from a modeling perspective, whether internally non-uniform groups strictly require the mixing prototypes mechanism and whether the latter unavoidably induces the former. On one hand, it appears conceptually very hard to implement the mixing mechanism in a manner that induces uniform groups or that does not induce any groups at all—at least without introducing highly arbitrary/implausible assumptions. On the other hand, it appears possible to define alternative procedures that are capable of generating non-uniform groups without making (explicit) use of mixing. Exploring such alternatives and their empirical

validity is beyond the purpose of this study, which did not aim at further empirical validation of the mixing ingredient, but just at using it as an easy, accessible way (due to availability from previous research) of generating internally non-uniform groups.

The fact that this study used multidimensional sociological data, while heavily relying on eigenvalue decomposition, may raise the question of how the approach here is different from traditional social science research using principal component analysis [37]. Although principal component analysis heavily relies on eigenvalue decomposition, in a social science context, the former most often implies a decomposition of the matrix of covariances or correlations between the variables, while this study focuses on the matrix of similarities between individuals. This actually makes the approach here conceptually more similar to clustering methods [38], which aim at identifying group structure, while providing an optimal partition of the given set of individuals. However, these methods do not attempt to decompose the similarity matrix and remove the irrelevant eigenmodes. In fact, following the approach of reference [14], the sum of the similarity matrix contributions associated to the structural modes identified here can be interpreted as a (modified) modularity matrix, which could provide a new method for clustering individuals via maximization of what one may call ‘spectral modularity’. Since this automatically eliminates the noise components and the common trend encoded in the global mode, such a method should be able to disentangle clusters that are not recognized by traditional approaches. However, such a method might also be sensitive to false positive cluster splittings, due to structural modes possibly being artifacts of feature redundancies, as shown in this study (at this point, it is not clear whether this is also a problem for the method in reference [14], intended for matrices of correlations between time series). There is certainly a need of finding new, eigenmode-dependent criteria for distinguishing feature redundancy artifacts from authentic groups that are internally non-uniform, criteria that would be generic enough to be used with empirical data—this would be valuable regardless of whether or not the structural modes are used for spectral modularity maximization. For such purposes, one could for instance speculate that eigenvectors associated to authentic groups should exhibit some degree of stability across different selections of features used for computing the similarity matrix, which might not be the case for eigenvectors associated to redundancies. Such effects would be first verified in simplified settings (like the one used here with the FCI, S2G and MPG models) and then potentially exploited in empirical contexts. These are some of the aspects left for future research.

8. Conclusion

This study examined cultural structure from a new angle, relying on notions from random matrix theory. It essentially provides a filtering procedure for matrices of cultural similarity between individuals, which eliminates, in a statistically robust way, the structurally-irrelevant components. Much effort was dedicated to the interpretation of the remaining, structurally-relevant components. On one hand, these may be a consequence of redundancies between cultural variables, mainly encoding information about the experimental setup. On the other hand, they may be a consequence of a modular organization of culture, thus encoding information about cultural groups. We have shown that it makes sense to conceptually distinguish between a ‘redundancies scenario’ and a ‘groups scenario’, as well as between internally uniform and internally non-uniform groups, even when cultural variables take very few, discrete values.

In empirical contexts, we were able to exclude internally uniform groups as a structural mechanism, but we were not able to exclude nor confirm internally non-uniform groups, which currently cannot be distinguished from redundancies. Discriminating between redundancies and internally non-uniform groups in generic empirical contexts requires further methodological research. This would allow, on one hand, to reject or accept the idea that culture has a modular structure, on the other hand to extend the applicability and reliability of the approach explored here, for the purpose of identifying subtle groups in discrete, multivariate data.

When it comes to random matrix studies of complex systems, the study also provides two messages of caution. First, it is not always appropriate to conclude that mesoscopic, system-specific groups are present, upon observing deviating eigenvalues. Second, it can be problematic to validate the associated eigenmodes via complementary, eigenvector-based quantities, which may be inherently biased against certain types of groups.

Data availability

The data that support the findings of this study are available upon reasonable request from the authors.

Acknowledgments

The author acknowledges insightful discussions with Diego Garlaschelli, Assaf Almog, Marco Verweij, Michael Thompson, Maroussia Favre, Jason Roos, Pieter Schoonees, Santo Fortunato and Vincent Traag, as well as financial support from the Netherlands Organization for Scientific Research (NWO/OCW), in particular via Grant No. 314-99-400.

Appendix A. Choice of null model

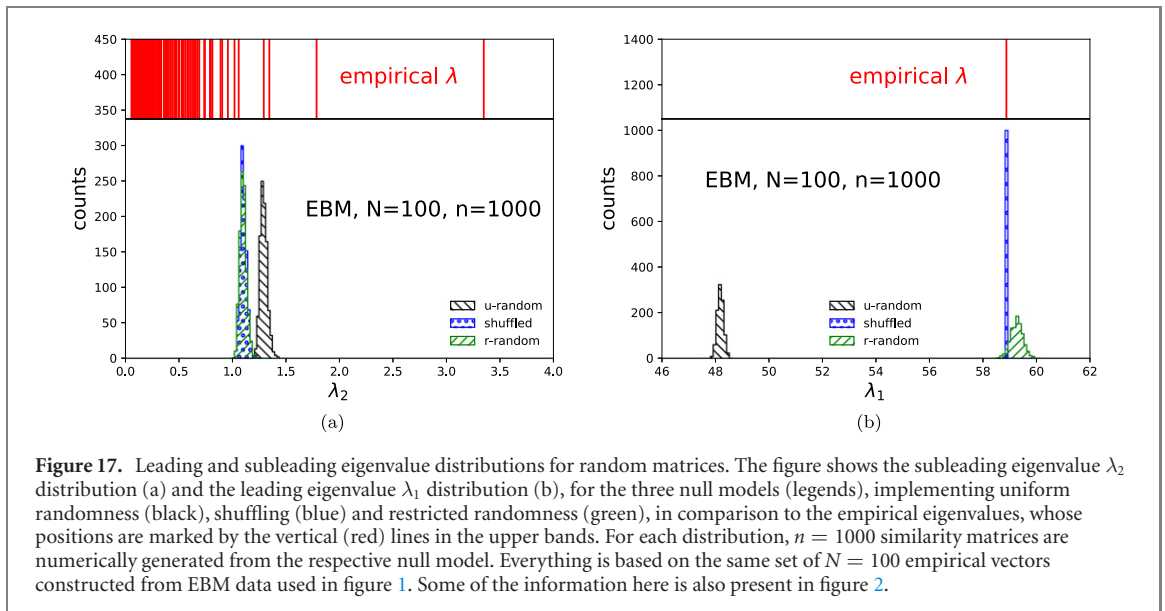
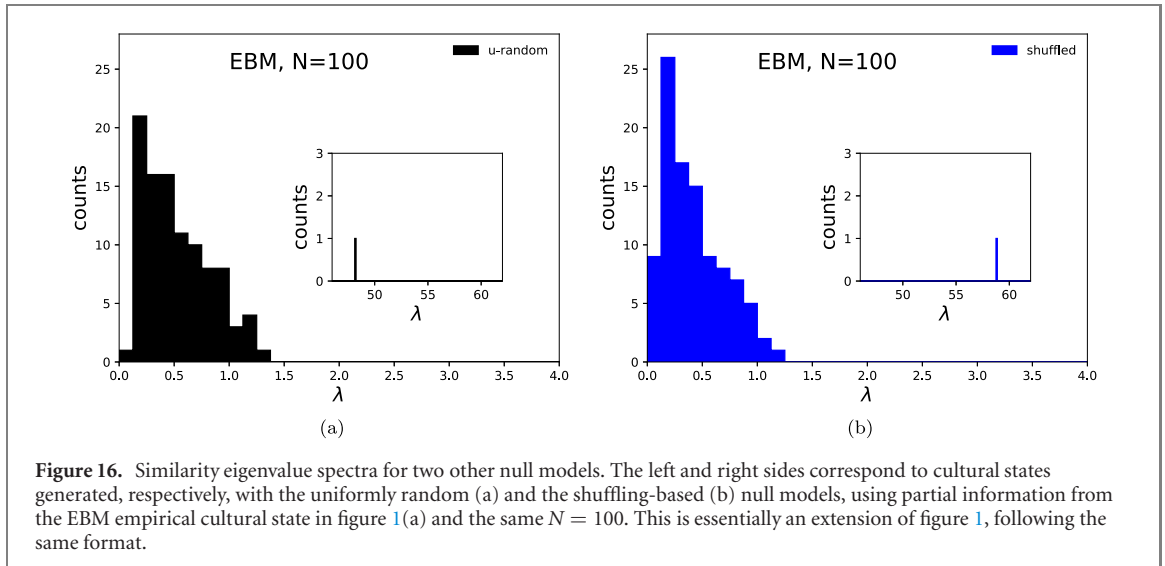
This section details the reasoning for choosing restricted randomness ('r-random') as the appropriate null model, intended for consistent use throughout this study, with particular relevance for section 2, section 5 and section 6. Comparisons are being made between restricted randomness and two other null models, based on uniform randomness and shuffling, often denoted as 'u-random' and 'shuffled' respectively. These comparisons focus on how well they reproduce the upper boundary of the noisy spectral region ('the bulk'), as well as the position of the leading eigenvalue, which for similarity matrices has a guaranteed [22] interpretation as the 'global mode'.

The uniformly random generation, already used in reference [22], is the simplest and arguably the most direct, similarity-oriented correspondent of the Marchenko–Pastur approach. Figure 16(a) shows the spectrum of a similarity matrix generated via uniform randomness, based on the empirical EBM state used in section 2. Specifically, for every vector, each trait is chosen independently at random from the traits available at the level of the respective feature, with the same probability attached to each possible trait. This means that uniform randomness retains minimal information from the empirical cultural state: only the number of features, the type and the number of traits of each feature. Note that the leading eigenvalue of this matrix is comparable to that of the empirical matrix, shown by figure 1(a). Reference [22] showed that the analytic, limiting distribution given by the Marchenko–Pastur formula has a shape that is qualitatively similar to the bulk of the u-random spectrum. Quantitatively however, the analytic and numerical distributions become truly similar only if an important parameter controlling the former is left free and fit to the numerical results, instead of being directly set to F/N , which can be done when dealing with matrices of correlations between N time series with F (numeric) entries each. Moreover, the Marchenko–Pastur formula completely fails to describe the leading eigenvalue.

A more sophisticated null model makes use of trait shuffling, which is known to be important for understanding empirical cultural states, independently from spectral decomposition and random matrix notions [5–9], since it reproduces exactly the empirical trait occurrence frequencies. Figure 16(b) shows the eigenvalue spectrum of a similarity matrix generated via shuffling, based on the empirical EBM state used in section 2. Specifically, with respect to every feature, the traits realized in the empirical state are randomly permuted among the vectors, such that every permutation is equally likely. This is done independently for every feature, so that all types of correlations between features are destroyed. The procedure preserves exactly the number of times each trait is empirically realized, in addition to preserving the data format of the empirical state. Note that the leading eigenvalue of the resulting matrix appears identical to that of the empirical matrix, shown by figure 1(a), a direct consequence of enforcing feature-level non-uniformities, which contribute to the global tendency toward similarity. This is the strength of the shuffled null model, compared to the u-random one. The weakness of shuffling consists of the assignment of traits to vectors being not entirely independent across vectors, implying that the number of vectors N resulting from shuffling has to be exactly the same as the number of empirical vectors used, which is not the case for uniform randomness.

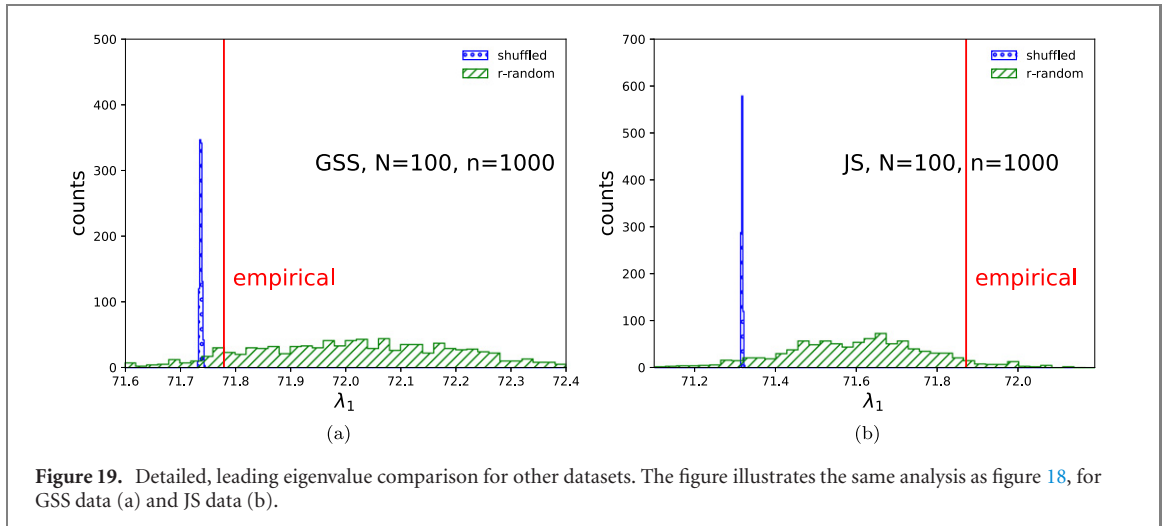
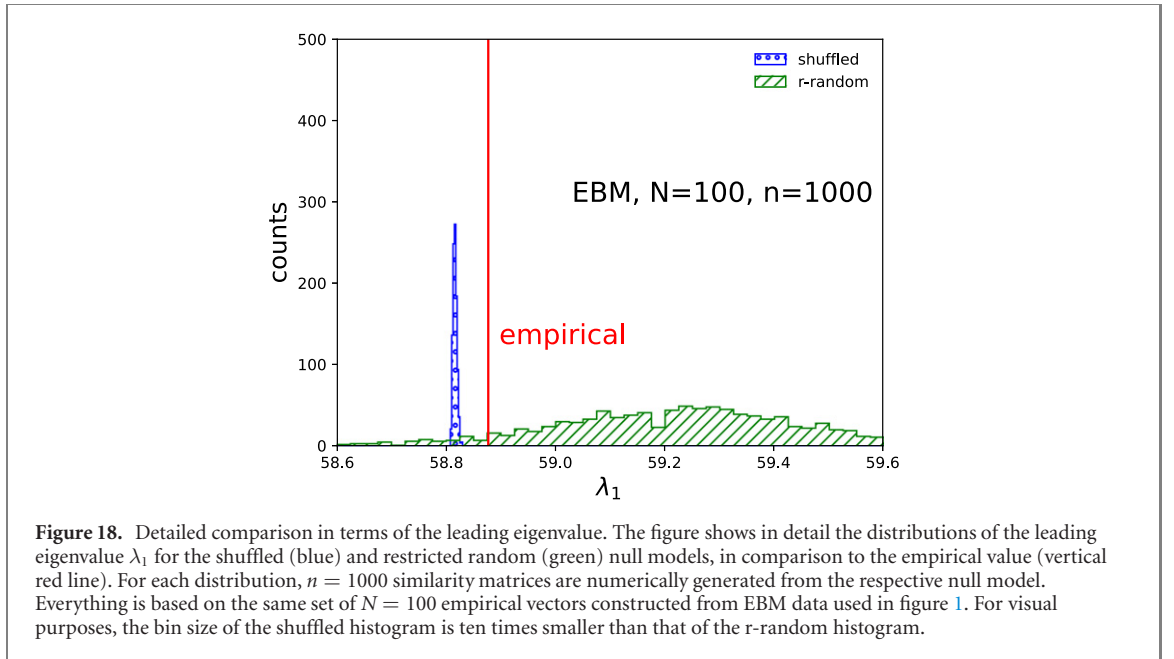
The r-random null model, introduced in this study, also reproduces, although only on average, the empirical trait frequencies, thus incorporating the strength of the shuffled model. Moreover, it also uses independent generation, like uniform randomness, allowing for an arbitrary number N of cultural vectors to be generated, thus avoiding the weakness of the shuffled model. The independent generation should also make the analytic tractability of the model easier than for shuffling. Although neither of these advantages is directly exploited in this study, they suggest that restricted randomness is conceptually more appropriate than either uniform randomness or shuffling, as it incorporates the desirable properties of both. The eigenvalue spectrum of an r-random similarity matrix is already shown in figure 1(b). Note that the shuffled and r-random cases reproduce much better the empirical leading eigenvalue than the u-random case. A similar point can be made with respect to reproducing the empirical bulk shape, after visually inspecting the lower- λ parts of the four histograms in figures 16 and 1.

Based on figures 1 and 16, one can already recognize that different null models can provide different predictions for the upper boundary of the noise-dominated spectral region, potentially leading to different sets of eigenpairs that are selected as structural modes. These differences between null models are directly related



to previously mentioned differences in terms of the leading eigenvalue, due to the eigenvalue summation constraint $\sum_l \lambda_l = N$ (see section 2). A better understanding of these aspects is provided by figure 17, which overcomes the limitations due to the binning and to the single matrix sampling inherent in figures 1 and 16. Specifically, figure 17(a) shows the subleading eigenvalue distribution for the three null models, while figure 17(b) shows the associated leading eigenvalue distributions. For comparison, the empirical eigenvalues are shown by the vertical (red) lines in the upper bands of figure 17. Each pair of λ_2 and λ_1 distributions is produced numerically, by sampling $n = 1000$ sets of cultural vectors from the statistical ensemble of the respective null model. Note that the shuffled and r-random λ_2 distributions are almost identical, while the u-random one is located at higher values. The latter would thus lead to the selection of only two structural modes (λ_2 and λ_3), while the former would lead to four structural modes (λ_2 to λ_5) that appear statistically significant, with departures of at least two standard deviations from the means of the respective distributions. Also note that the shuffled and r-random λ_1 distributions, although different in shape, are also very close to each other and to the empirical value, while the u-random λ_1 distribution corresponds to significantly lower values. This is in agreement with previous observations, as well as with the idea that, when exploring different null models, lower λ_1 correspond to higher λ_2 , due to the eigenvalue summation constraint. Moreover, it shows more clearly that the overall tendency toward similarity is smaller in the u-random case than in the empirical, shuffled and r-random cases, since the last 3 cases share the same feature-level probability profiles, which are generally not uniform, unlike the u-random case.

In figure 17(b), the empirical λ_1 also appears statistically compatible with both shuffling and restricted randomness, but closer to the mean of the former. This, however, deserves a closer inspection, due to the



limitations inherent in the binning of figure 17(b). Figure 18 focuses on the shuffled and r-random λ_1 distributions, giving a better impression of how well either null model predicts the empirical leading eigenvalue based on partial information about trait frequencies. It appears that, due to the sharpness of the shuffled λ_1 distribution, the empirical value is actually not statistically compatible with it, while it is clearly compatible with the r-random distribution. This is consistent with the fact that trait shuffling is inherently more rigid than restricted randomness, in a probabilistic and information theoretic sense (the former entirely forbids certain cultural vectors, while the latter only attaches relatively low probabilities to them), while constituting another advantage of restricted randomness over trait shuffling. In order to further validate this idea, the analysis in figure 18 is repeated for the GSS and JS empirical states introduced in section 2. The results are shown in figure 19: restricted randomness indeed provides a better description of the empirical global mode also in these cases. Moreover, in the JS case (figure 19(b)), the empirical value is closer to the r-random distribution than to the shuffled distribution also in terms of the difference from the mean. Thus, any idea that estimates based on trait shuffling might still be closer to reality (though systematically biased), that might be suggested by the EBM (figure 18) and GSS (figure 19(a)) cases, is not valid in general.

Considering all the above, we choose restricted randomness as the appropriate null model for use throughout the manuscript. Besides being most favorable on theoretical, *a priori* grounds (conveniently combining the independent sampling of uniform randomness with the trait frequency replication of shuffling), it is also best behaved in a computational, *a posteriori* sense (providing the best predictions for the leading eigenvalue, which is intimately coupled to the random bulk boundary). The choice provides a filtering procedure that allows the researcher to identify empirical structure that is present independently of feature-level non-uniformities,

which are expected to strongly depend on how the associated items and the possible values are formulated and much less on authentic properties of the real social system from which the data is extracted.

Appendix B. Proof of positive semidefiniteness

This section presents a proof of the fact that cultural similarity matrices of a ‘Hamming–Manhattan’ type, as defined via equation (1), are all positive semidefinite. Mathematically, the statement may be written as: $S \geq 0$ for any real vector $\vec{w}$, where S is the scalar quantity defined in equation (C2). It is of great use to first reformulate the statement in terms of feature-level similarities. Specifically, by inserting equation (1) in (C2) one finds that:

$$S = \frac{1}{N} \sum_{k=1}^F \left(\sum_{i=1}^N \sum_{j=1}^N w_i s_{ij}^k w_j \right), \quad (\text{B1})$$

with:

$$s_{ij}^k = f_{\text{nom}}^k \delta(x_i^k, x_j^k) + (1 - f_{\text{nom}}^k) \left(1 - \frac{|x_i^k - x_j^k|}{q^k - 1} \right) \quad (\text{B2})$$

denoting one of the F single-feature similarity matrices. It becomes clear that any aggregate similarity matrix $(s_{ij})_{i,j \in \overline{1,N}}$ is positive semidefinite if and only if any single-feature matrix $(s_{ij}^k)_{i,j \in \overline{1,N}}$ is positive semidefinite. The sufficiency of the latter condition results from a simple averaging over feature level inequalities of the form: $\vec{w}^T s^k \vec{w} \geq 0$ to obtain $S \geq 0$, based on equation (B1). On the other hand, the necessity results from the observation that, in general, an aggregate similarity matrix may use any feature configuration, including any single-feature configuration.

To complete the proof, we show that any single-feature similarity matrix is positive semidefinite, first for the nominal (Hamming) case in appendix B.1, second for the ordinal (Manhattan) case in appendix B.2. These two sections make exclusive use of the ‘ s_{ij} ’ notation for entries of single-feature similarity matrices, instead of the ‘ s_{ij}^k ’ notation. This simplification is also reflected by the use of ‘ x_i ’ instead of ‘ x_i^k ’ when denoting the entry of cultural vector i for the respective feature.

B.1. Nominal single-feature similarity

In order to prove that a similarity matrix $(s_{ij})_{i,j \in \overline{1,N}}$ constructed from one, nominal feature is positive semidefinite, we show that $S \geq 0$ for any real vector $\vec{w}$, where S is the scalar quantity defined in equation (C2). This translates to:

$$\sum_{i=1}^N w_i^2 + 2 \sum_{i=1}^{N-1} \sum_{j=i+1}^N w_i w_j \delta(x_i, x_j) \geq 0, \quad (\text{B3})$$

after having used the fact that $s_{ii} = 1$ and that $s_{ij} = \delta(x_i, x_j)$, resulting from equation (B2).

It is important to note that the nominal feature induces a clustering (or partition) $\mathcal{O}$ of the set of vectors $\overline{1,N}$, so that all vectors picking a specific trait belong to a certain cluster (or part) O . Using this observation, together with the definition of the δ -function, one may rewrite equation (B3) as:

$$\sum_{O \in \mathcal{O}} \left[\sum_{i \in O} w_i^2 + \sum_{i \in O} \sum_{j \in O - \{i\}} w_i w_j \right] \geq 0, \quad (\text{B4})$$

which can be further reduced to:

$$\sum_{O \in \mathcal{O}} \left[\left(\sum_{i \in O} w_i \right)^2 \right] \geq 0, \quad (\text{B5})$$

which is a valid statement, since the left-hand side is a sum of non-negative numbers. This concludes the proof for the nominal (Hamming), single-feature case.

B.2. Ordinal single-feature similarity

In order to prove that a similarity matrix $(s_{ij})_{i,j \in \overline{1,N}}$ constructed from one, ordinal feature is positive semidefinite, we show that an extension of Sylvester’s criterion, namely Prussing’s criterion [39] is satisfied: that all principal minors are non-negative. First, note that all principal minors of order 1 correspond to the diagonal elements and are thus equal to 1. Thus, the proof focuses on higher order principal minors. These are essentially determinants of smaller similarity matrices, based on the same ordinal feature and on subsets of cultural vectors sampled from the larger set associated to the larger matrix. Thus, the proof reduces to showing that the determinant of any ordinal, single-feature similarity matrix with $N \geq 2$ elements is non-negative.

Based on equation (B2), such a similarity matrix can be written as:

$$s_{ij} = 1 - \frac{|x_i - x_j|}{q-1}, \quad (\text{B6})$$

which makes it clear that the entire matrix is specified by the relative positioning of N rational numbers $x_i/(q-1)$ within the $[0, 1]$ interval. Let $x_{i'}/(q-1)$ denote the same rational numbers but sorted for increasing values, so that: $x_{i'} \leq x_{i'+1}$, $\forall i' \in \overline{1, N-1}$. This amounts to permuting the rows and columns of s in the same way, leaving the value of its determinant $\det(s)$ unchanged. After this operation, the determinant may be conveniently written in terms of the (Manhattan) distance values $d_{i'} = 1 - s_{i'(i'+1)}$ associated to pairs of consecutive numbers, in the following way:

$$\det(s) = \begin{vmatrix} 1 & 1-d_1 & 1-d_1-d_2 & \dots \\ 1-d_1 & 1 & 1-d_2 & \dots \\ 1-d_1-d_2 & 1-d_2 & 1 & \dots \\ \dots & \dots & \dots & \dots \end{vmatrix}. \quad (\text{B7})$$

This can be brought to an upper-diagonal form by applying further row and column operations that conserve the determinant (up to a minus sign). Specifically, the following steps are taken:

- Subtract row $i+1$ from row i for every $i \in \overline{1, N-1}$
- Add column N to every column $j \in \overline{1, N-1}$
- Bring row N to the top (by exchanging it with rows $N-1$ to 1), while producing a factor of $(-1)^{N-1}$ to obtain:

$$\det(s) = (-1)^{N-1} \begin{vmatrix} 2 - \sum_{l=1}^{N-1} d_l & 2 - \sum_{l=2}^{N-1} d_l & 2 - \sum_{l=3}^{N-1} d_l & \dots & 2 - d_{N-1} & 1 \\ 0 & -2d_1 & -2d_1 & \dots & -2d_1 & -d_1 \\ 0 & 0 & -2d_2 & \dots & -2d_2 & -d_2 \\ \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & 0 & \dots & -2d_{N-2} & -d_{N-2} \\ 0 & 0 & 0 & \dots & 0 & -d_{N-1} \end{vmatrix}, \quad (\text{B8})$$

which evaluates to the product of the diagonal elements:

$$\det(s) = 2^{N-2} \left(2 - \sum_{i=1}^{N-1} d_i \right) \prod_{j=1}^{N-1} d_j, \quad (\text{B9})$$

where $d_i \geq 0$, $\forall i \in \overline{1, N}$ (non-negativity of Manhattan distances) and $\sum_{i=1}^{N-1} d_i \leq 1$ (the sum is equal to the difference between the highest and the lowest of the $x_i/(q-1)$ values, which are constrained to the $[0, 1]$ interval). Thus, $\det(s) \geq 0$, which concludes the proof for the ordinal (Manhattan), single-feature case.

Appendix C. The mathematics of similarity eigenpairs

This section presents some elementary, but important mathematical properties of the eigenvalues λ_l and the associated eigenvectors v_l satisfying equation (2). For the sake of clarity, the following explanations make use of the term ‘individual’ as a replacement for ‘cultural vector’, although most of the concepts presented are also valid, at least mathematically, for similarity matrices constructed from randomly generated cultural vectors, based on any probabilistic model.

Since the eigenvectors v_l have only real entries and form an orthonormal basis, one can write any real vector w with N entries as a linear combinations of the eigenvectors:

$$w = \sum_{l=1}^N \alpha_l v_l, \quad (\text{C1})$$

with real coefficients α_l . The rest of this argument is restricted to unit vectors w , which satisfy $\sum_{i=1}^N w_i^2 = 1$, which can be translated as $\sum_{l=1}^N \alpha_l^2 = 1$ in terms of the eigenvectors’ coefficients. This encompasses all the eigenvectors $w = v_l, \forall l$ as special cases. Moreover, let us define the following scalar quantity:

$$S = \sum_{i=1}^N \sum_{j=1}^N w_i s_{ij} w_j, \quad (\text{C2})$$

as the double contraction of the similarity matrix s with the vector w . By means of equations (C2) and (C1), for any vector w (including the special cases when this entirely matches one of the eigenvectors v_l) every entry of w becomes associated to one of the individuals based on which the similarity matrix s is computed. Thus, w can be seen as a (normalized) linear combination of the N individuals. S can be then interpreted as the self-similarity of any normalized linear combination w , since every pairwise similarity s_{ij} is multiplied by the numbers w_i and w_j attached to individuals i and j . For any normalized w , one can show that:

$$S = 1 + 2 \sum_{i=1}^{N-1} \sum_{j=i+1}^N w_i s_{ij} w_j, \quad (\text{C3})$$

which immediately follows from the fact that $s_{ii} = 1, \forall i$, which is a direct consequence of how the similarity is defined in equation (1). Note that $S = 1$ whenever w gives a strength of 1 to one individual and 0 to all the others, which supports the interpretation of S as a self similarity. It is also important to note, from equation (C3), that S is larger when w is such that pairs of entries (i, j) with the same sign correspond to higher values of s_{ij} and higher values of $|w_i w_j|$, while pairs with opposite signs correspond to lower values of s_{ij} and lower values of $|w_i w_j|$.

The largest self-similarity S is attained when the linear combination w , among all unit vectors, takes the form of the eigenvector v_1 with the largest associated eigenvalue λ_1 , corresponding to $\alpha_l = \delta_l^1, \forall l$. This largest self-similarity value is actually equal to the largest eigenvalue: $S = \lambda_1$. This is shown by plugging equations (2) and (C1) into (C2) and using the normalization condition, leading to:

$$S = \sum_{l=1}^N \alpha_l^2 \lambda_l. \quad (\text{C4})$$

More generally, one can see here that each eigenvector v_l with the l th highest eigenvalue λ_l , corresponding to $\alpha_{l'} = \delta_{l'}^l, \forall l'$, is such that it gives the largest possible value of $S = \lambda_l$, while also being normalized and orthogonal to all eigenvectors $v_{l'}$ with $\lambda_{l'} > \lambda_l$. When confronting this with the insights provided by equation (C3), one realizes that any subset of individuals with strong, internal similarities is captured by one of the eigenmodes, whose eigenvalue is larger if the overall level of internal similarity is higher. Moreover, the eigenvector entries of these strongly similar elements will have the same sign and the highest absolute values.

Appendix D. The FCI model

This section gives the details behind the mathematical expressions in section 3.1, which introduced the FCI model. Deriving the probability distribution p associated to this model follows the maximum-entropy approach introduced by reference [26]. This crucially relies on the Shannon entropy, which is a functional of the probability distribution:

$$H[p] = - \sum_{\vec{s}} p_{\vec{s}} \log p_{\vec{s}}, \quad (\text{D1})$$

where $\vec{s}$ denotes a generic spin configuration with F spins on a fully-connected lattice, or a generic cultural vector with F binary cultural features whose possible traits are marked as ‘−1’ and ‘+1’. The value of the functional H is maximized subject to two constraints, one related to the normalization of the probability distribution over the set of possible configurations:

$$\sum_{\vec{s}} p_{\vec{s}} = 1, \quad (\text{D2})$$

the other related to enforcing, on average, a certain amount K of alignment:

$$\sum_{a < b} \sum_{\vec{s}} S_a S_b p_{\vec{s}} = K, \quad (\text{D3})$$

namely the average number of pairs of similarly labeled traits within a given configuration $\vec{s}$, where the first summation is over all distinct pairs of distinct features (or lattice sites). The maximization is done using the Lagrange multipliers technique for equations (D1)–(D3), which implies that one should find the extrema of the following functional:

$$L[p] = H[p] - \lambda_0 \left(\sum_{\vec{s}} p_{\vec{s}} - 1 \right) - \lambda \left(\sum_{a < b} \sum_{\vec{s}} S_a S_b p_{\vec{s}} - K \right), \quad (\text{D4})$$

where λ_0 and λ are free parameters associated to the two constraints. By taking partial derivatives of equation (D4) with respect to each $p_{\vec{s}}$ and further manipulations, one finds the following probability distribution:

$$p_{\vec{s}} = \frac{1}{Z(-\lambda)} \exp \left[-\lambda \sum_{a < b} S_a S_b \right], \quad (\text{D5})$$

where $Z(-\lambda)$ is a normalization factor, known in statistical physics as the ‘partition function’:

$$Z(-\lambda) = \sum_{\vec{s}} \exp \left[-\lambda \sum_{a < b} S_a S_b \right], \quad (\text{D6})$$

where one can replace the coupling parameter $-\lambda$ with $\mu > 0$ (whose positive value favors alignment as opposed to anti-alignment, which corresponds to ferromagnetism) and re-express the sum over configurations $\vec{s}$ as a sequence of sums over the possible traits of each feature S_k , leading to:

$$Z(\mu) = \prod_{k=1}^F \left(\sum_{S_k=\pm 1} \right) \exp \left[\mu \sum_{a=1}^{F-1} \sum_{b=a+1}^F S_a S_b \right]. \quad (\text{D7})$$

In the exponent of this expression, there are $F(F-1)/2$ terms, out of which $F_+(F-F_+)$ are equal to -1 , while the other are equal to $+1$. Based on this, after further manipulations and after taking advantage of symmetries, the partition function can be expressed as:

$$Z(\mu) = \sum_{F_+=0}^F \frac{F!}{F_+!(F-F_+)!} \exp \left[\frac{\mu}{2} ((2F_+ - F)^2 - F) \right], \quad (\text{D8})$$

where the combinatorial factor (binomial coefficient) before the exponential function counts the number of configurations with the same number F_+ of $+1$ traits. In a way rather analogous to the partition function, the double summation in the exponent of equation (D5) can also be eliminated. After multiplication with the combinatorial factor, this leads to equation (3), which gives the probability of having a configuration with F_+ spins up.

On the other hand, using equations (D6), (D3) can be written as:

$$K = -\frac{\partial(\log(Z(-\lambda)))}{\partial \lambda} = \frac{\partial(\log(Z(\mu)))}{\partial \mu}, \quad (\text{D9})$$

while the correlation between features/spins a and b is:

$$C_{ab} = \frac{\langle S_a S_b \rangle - \langle S_a \rangle \langle S_b \rangle}{\sqrt{\langle S_a^2 \rangle - \langle S_a \rangle^2} \sqrt{\langle S_b^2 \rangle - \langle S_b \rangle^2}}, \quad (\text{D10})$$

where $\langle Q \rangle = \sum_{\vec{s}} Q_{\vec{s}} p_{\vec{s}}$ is the expected value of quantity Q with respect to the statistical ensemble. One can show—using symmetry arguments or equations (D5)–(D7)—that $\langle S_a^2 \rangle = 1$ and that $\langle S_a \rangle = \langle S_b \rangle = 0$, so $C_{ab} = \langle S_a S_b \rangle = \sum_{\vec{s}} S_a S_b p_{\vec{s}}$, which combined with equation (D3) leads to $\sum_{a < b} C_{ab} = K$. But due to symmetry, the expected correlation C_{ab} is the same for all pairs (a, b) , so:

$$C_{ab} = C(\mu, F) = \frac{2}{F(F-1)} K = \frac{2}{F(F-1)} \frac{\partial(\log(Z(\mu)))}{\partial \mu}, \quad (\text{D11})$$

for any pair (a, b) , which (after making use of equation (D8)) can also be written in the form shown by equation (4)—equation (D9) was used for the last transformation in equation (D11).

One should expect that $C(0.0) = 0.0$ (null correlations for null coupling), which based on equation (4), implies that the following identity holds:

$$\sum_{F_+=0}^F \frac{(F-2)! ((2F_+ - F)^2 - F)}{F_+!(F-F_+)!} = 0, \quad (\text{D12})$$

which, after substitution of F_+ with k and of F with N and some further manipulations leads to the following combinatorial identity:

$$\sum_{k=0}^N \binom{N}{k} ((2k - N)^2 - N) = 0, \quad (\text{D13})$$

which can be shown to hold using the expressions for the binomial expansion and for the first and second moments of a binomial distribution with the probability parameter set to 0.5.

Appendix E. The S2G model

This section provides the mathematical derivations of the important mathematical formulas related to the S2G model, introduced in section 3.2. The derivations are based on the model description there.

First, we prove equation (5). On one hand, the probability that a cultural vector meant to be part of group +1 is assigned to a configuration with F_+ traits +1 is:

$$p_+^+(\nu, F, F_+) = \frac{F!}{F_+!(F - F_+)!} (1 - 2\nu)^{F_+} (2\nu)^{F - F_+}, \quad (\text{E1})$$

which is a binomial distribution with probability $1 - 2\nu$ for the +1 possibility and 2ν for the -1 possibility. On the other hand, the probability that a configuration meant to be part of group -1 has F_+ traits +1 is:

$$p_+^-(\nu, F, F_+) = \frac{F!}{F_+!(F - F_+)!} (2\nu)^{F_+} (1 - 2\nu)^{F - F_+}, \quad (\text{E2})$$

which is the same binomial distribution, but with inverted probabilities. Since the two groups are by construction equally likely, the combined probability of all configurations with F_+ traits +1 is:

$$p(\nu, F, F_+) = \frac{1}{2} p_+^+(\nu, F, F_+) + \frac{1}{2} p_+^-(\nu, F, F_+). \quad (\text{E3})$$

Inserting equations (E1) and (E2) in (E3) leads to equation (5).

Second, we prove equation (6). The correlation coefficient of any two features a and b is given by equation (D10), which, for symmetry reasons similar to the case of the FCI model, simplifies to:

$$C_{ab}(\nu) = \sum_{\vec{S}} S_a S_b p_{\vec{S}}(\nu) = C(\nu). \quad (\text{E4})$$

Moreover, the probability attached to any configuration $\vec{S}$ can be written as:

$$p_{\vec{S}}(\nu) = \frac{1}{2} \left(p_{\vec{S}}^-(\nu) + p_{\vec{S}}^+(\nu) \right), \quad (\text{E5})$$

where $p_{\vec{S}}^-(\nu)$ and $p_{\vec{S}}^+(\nu)$ are the probabilities of configuration $\vec{S}$, conditional on whether it is generated for group -1 or for group +1 respectively. In turn, these probabilities can be factorized in terms of feature-level probabilities of traits:

$$p_{\vec{S}}^-(\nu) = \prod_{a=1}^F p_{S_a}^-(\nu), \quad p_{\vec{S}}^+(\nu) = \prod_{a=1}^F p_{S_a}^+(\nu), \quad (\text{E6})$$

because once the group is chosen, each trait S_a (with possible values -1 and +1) is chosen independently at the level of the respective feature a . By inserting equation (E6) in (E5) and the result in equation (E4), by carrying out appropriate algebraic manipulations, while making use of the fact that $\sum_{\vec{S}} = \prod_{a=1}^F (\sum_{S_a})$ and of the fact that $p_{S_a=-1}^{+/+}(\nu) + p_{S_a=+1}^{+/+}(\nu) = 1.0$, one obtains:

$$C(\nu) = \frac{1}{2} [p_{--}^-(\nu) - p_{-+}^-(\nu) - p_{+-}^-(\nu) + p_{++}^-(\nu)] + \frac{1}{2} [p_{--}^+(\nu) - p_{-+}^+(\nu) - p_{+-}^+(\nu) + p_{++}^+(\nu)], \quad (\text{E7})$$

where, for instance, $p_{-+}^-(\nu)$ is the probability that trait -1 is chosen for one of the features and that trait +1 is chosen for the other feature, conditional on the given configuration being generated for group -1. Based on the model description in section 3.2, one can see that:

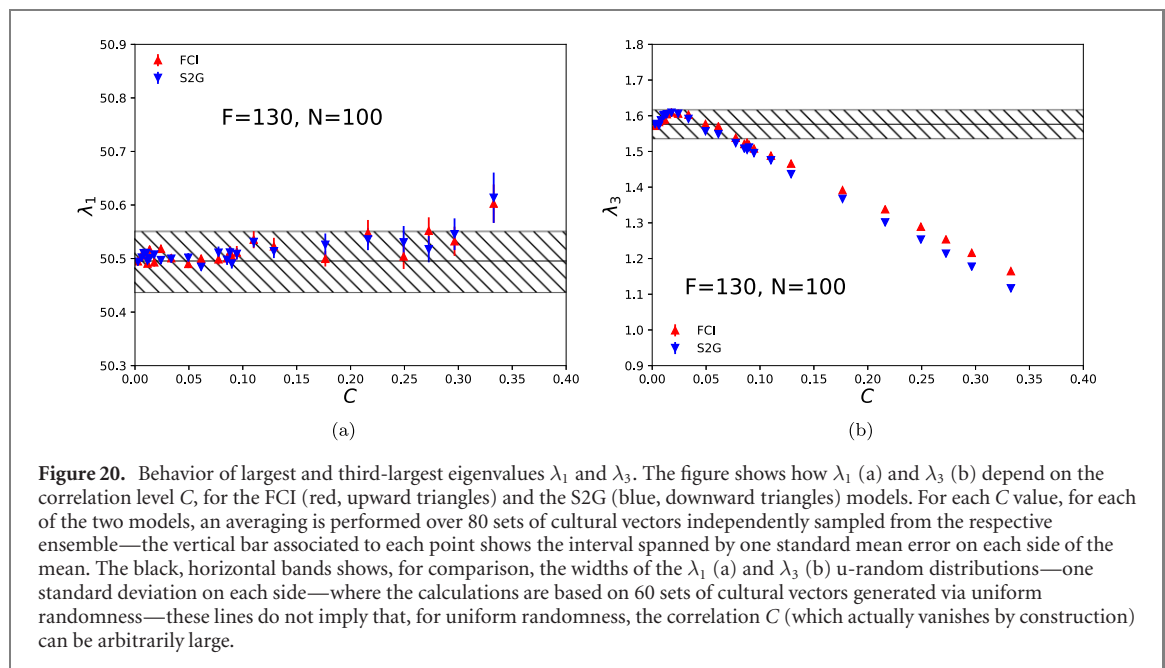
$$p_{--}^-(\nu) = p_{++}^+(\nu) = (1 - 2\nu)^2, \quad (\text{E8})$$

$$p_{-+}^-(\nu) = p_{-+}^+(\nu) = (2\nu)^2, \quad (\text{E9})$$

$$p_{+-}^-(\nu) = p_{+-}^+(\nu) = (1 - 2\nu)(2\nu), \quad (\text{E10})$$

$$p_{++}^-(\nu) = p_{++}^+(\nu) = (2\nu)(1 - 2\nu). \quad (\text{E11})$$

By plugging these in equation (E7), after simple algebraic manipulations, one obtains equation (6).



Appendix F. The structure of the FCI and S2G models

This section shows that the structure implicit in cultural states generated with either the FCI or the S2G model is mostly captured by only one eigenpair of the similarity matrix, so that there is at most one structural mode. Specifically, as the correlation level is increased for the FCI and the S2G models, there is only one eigenvalue—the subleading eigenvalue λ_2 —that becomes separated from the random bulk, while becoming significantly larger than the upper boundary of the bulk that is expected based on uniform randomness. The behavior of λ_2 has already been presented in figure 6. The results shown here, via figure 20, are complementary to those shown in figure 6, while using the same format. Specifically, figure 20(a) shows the behavior of λ_1 , while figure 20(b) shows the behavior of λ_3 . Note that λ_1 , associated to the global mode, remains statistically compatible with the null model as the level of correlation is increased, for both FCI and S2G. On the other hand, λ_3 decreases, while becoming, for large enough C , significantly smaller than the upper boundary of the bulk predicted by uniform randomness. All this shows that the structure FCI and S2G is mostly captured by the eigenpair of λ_2 , which becomes increasingly stronger as the correlation level increases. This appears to be a consequence of the fact that each model is controlled by one parameter, while all the non-uniformity of the associate probability distribution is captured by one dimension, namely the F_+ axis of figure 4.

ORCID iDs

Alexandru-Ionuț Băbeanu  <https://orcid.org/0000-0001-6274-4871>

References

- [1] Urry J 2005 The complexity turn *Theor. Cult. Soc.* **22** 1–14
- [2] Lazer D *et al* 2009 Social science: computational social science *Science* **323** 721–3
- [3] Kadushin C 2012 *Understanding Social Networks: Theories, Concepts and Findings* (Oxford: Oxford University Press)
- [4] Castellano C, Fortunato S and Loreto V 2009 Statistical physics of social dynamics *Rev. Mod. Phys.* **81** 591–646
- [5] Valori L, Picciolo F, Allansdottir A and Garlaschelli D 2012 Reconciling long-term cultural diversity and short-term collective social behavior *Proc. Natl Acad. Sci.* **109** 1068–73
- [6] Stivala A, Robins G, Kashima Y and Kirley M 2014 Ultrametric distribution of culture vectors in an extended Axelrod model of cultural dissemination *Sci. Rep.* **4** 4870
- [7] Băbeanu A-I, Talman L and Garlaschelli D 2017 Signs of universality in the structure of culture *Eur. Phys. J. B* **90** 237
- [8] Băbeanu A-I and Garlaschelli D 2018 Evidence for mixed rationalities in preference formation *Complexity* **2018** 3615476
- [9] Băbeanu A-I, van de Vis J and Garlaschelli D 2018 Ultrametricity increases the predictability of cultural dynamics *New J. Phys.* **20** 103026
- [10] Axelrod R 1997 The dissemination of culture *J. Conflict Resolut.* **41** 203–26
- [11] Mehta M L 2012 *Random Matrices* (Amsterdam: Elsevier)
- [12] Edelman A and Rao N R 2005 Random matrix theory *Acta Numer.* **14** 233–97
- [13] Potters M, Bouchaud J-P and Laloux L 2005 Financial applications of random matrix theory: old laces and new pieces *Acta Phys. Pol. B* **36** 2767–84

- [14] MacMahon M and Garlaschelli D 2015 Community detection for correlation matrices *Phys. Rev. X* **5** 021006
- [15] Anagnostou I, Squartini T, Garlaschelli D and Kandhai D 2020 Uncovering the mesoscale structure of the credit default swap market to improve portfolio risk modelling (arXiv:2006.03014v1)
- [16] Ali N, Ardalankia J, Raei R, Hedayatifar L, Ali H, Haven E and Reza Jafari G 2020 Analysis of the global banking network by random matrix theory (arXiv:2007.14447v1)
- [17] Bhosale U T, Tekur S H and Santhanam M S 2018 Scaling in the eigenvalue fluctuations of correlation matrices *Phys. Rev. E* **98** 052133
- [18] Barucca P, Kieburg M and Ossipov A 2020 Eigenvalue and eigenvector statistics in time series analysis *Europhys. Lett.* **129** 60003
- [19] Almog A, Buijink M R, Roethler O, Michel S, Johanna H, Meijer J H, Rohling T and Garlaschelli D 2019 Uncovering functional signature in neural systems via random matrix theory *PLOS Comput. Biol.* **15** 1–20
- [20] Akiki T J and Abdallah C G 2019 Determining the hierarchical architecture of the human brain using subject-level clustering of functional networks *Sci. Rep.* **9** 19290
- [21] Marchenko V A and Pastur L A 1967 Distribution of eigenvalues for some sets of random matrices *Math. USSR-Sb.* **1** 457–83
- [22] Patil A and Santhanam M S 2015 Random matrix approach to categorical data analysis *Phys. Rev. E* **92** 032130
- [23] Reif K and Anna M 1995 Euro-barometer 38.1: consumer protection and perceptions of science and technology, November 1992 (<http://icpsr.umich.edu/icpsrweb/ICPSR/studies/06045>)
- [24] Smith T W, Marsden P, Hout M and Kim J 1972–2012 General social surveys, 1993 ed (<http://gss.norc.oregonstate.edu/get-the-data/spss>)
- [25] Goldberg K, Roeder T, Gupta D and Perkins C 2001 Eigentaste: a constant time collaborative filtering algorithm *Inf. Retr.* **4** 133–51
- [26] Jaynes E T 1957 Information theory and statistical mechanics *Phys. Rev.* **106** 620–30
- [27] Colonna-Romano L, Gould H and Klein W 2014 Anomalous mean-field behavior of the fully connected ising model *Phys. Rev. E* **90** 042111
- [28] Reichl L E 2016 *A Modern Course in Statistical Physics* 4th revised and updated edn (New York: Wiley)
- [29] Goldenfeld N 1992 *Lectures on Phase Transitions and the Renormalization Group* (Boca Raton, FL: CRC Press)
- [30] Solé R V, Manrubia S C, Luque B, Delgado J and Bascompte J 1996 Phase transitions and complex systems: simple, nonlinear models capture complex systems at the edge of chaos *Complexity* **1** 13–26
- [31] Jones K R W 1990 Entropy of random quantum states *J. Phys. A: Math. Gen.* **23** L1247
- [32] Chakraborti A, Hrishidev , Sharma K and Pharasi H K 2020 Phase separation and scaling in correlation structures of financial markets (arXiv:1910.06242v3)
- [33] Thompson M, Ellis R J and Wildavsky A 1990 *Cultural Theory* (Boulder, CO: Westview Press)
- [34] Zhao X, Shang P and Huang J 2017 Mutual-information matrix analysis for nonlinear interactions of multivariate time series *Nonlinear Dyn.* **88** 477–87
- [35] Shang D, Xu M and Shang P 2017 Generalized sample entropy analysis for traffic signals based on similarity measure *Physica A* **474** 1–7
- [36] He J and Shang P 2018 Multidimensional scaling analysis of financial stocks based on kronecker-delta dissimilarity *Commun. Nonlinear Sci. Numer. Simul.* **63** 186–201
- [37] Duntelman G H 1989 *Principal Components Analysis* (Newbury Park, CA: SAGE Publications)
- [38] Kaufman L and Rousseeuw P J 1990 *Finding Groups in Data: An Introduction to Cluster Analysis* (New York: Wiley)
- [39] Prussing J E 1986 The principal minor test for semidefinite matrices *J. Guid. Control Dyn.* **9** 121–2